

ELECTRICAL SCIENCE SERIES

*An Introduction
To The Theory Of
Microwave Circuits*

ACADEMIC PRESS. New York and London, 1969

K. Kurokawa

BELL TELEPHONE LABORATORIES, INC.



ELECTRICAL SCIENCE

A Series of Monographs and Texts

Edited by

Henry G. Booker

UNIVERSITY OF CALIFORNIA AT SAN DIEGO
LA JOLLA, CALIFORNIA

Nicholas DeClaris

UNIVERSITY OF MARYLAND
COLLEGE PARK, MARYLAND

- JOSEPH E. ROWE. Nonlinear Electron-Wave Interaction Phenomena. 1965
- MAX J. O. STRUTT. Semiconductor Devices: Volume I. Semiconductors and Semiconductor Diodes. 1966
- AUSTIN BLAQUIERE. Nonlinear System Analysis. 1966
- VICTOR RUMSEY. Frequency Independent Antennas. 1966
- CHARLES K. BIRDSALL AND WILLIAM B. BRIDGES. Electron Dynamics of Diode Regions. 1966
- A. D. KUZ'MIN AND A. E. SALOMONOVICH. Radioastronomical Methods of Antenna Measurements. 1966
- CHARLES COOK AND MARVIN BERNFELD. Radar Signals: An Introduction to Theory and Application. 1967
- J. W. CRISPIN, JR., AND K. M. SIEGEL (eds.). Methods of Radar Cross Section Analysis. 1968
- GIUSEPPE BIORCI (ed.). Network and Switching Theory. 1968
- ERNEST C. OKRESS (ed.). Microwave Power Engineering:
Volume 1. Generation, Transmission, Rectification. 1968
Volume 2. Applications. 1968
- T. R. BASHKOW (ed.). Engineering Applications of Digital Computers. 1968
- JULIUS T. TOU (ed.). Applied Automata Theory. 1968
- ROBERT LYON-CAEN. Diodes, Transistors, and Integrated Circuits for Switching Systems. 1969
- M. RONALD WOHLERS. Lumped and Distributed Passive Networks. 1969
- MICHEL CUENOD AND ALLEN E. DURLING. A Discrete-Time Approach for System Analysis. 1969
- K. KUROKAWA. An Introduction to the Theory of Microwave Circuits. 1969

In Preparation

- GEORGE TYRAS. Radiation and Propagation of Electromagnetic Waves.
- GEORGE METZGER AND JEAN PAUL VABRE. Transmission Lines with Pulse Excitation.
- C. L. SHENG. Threshold Logic.
- HUGO MESSERLE. Energy Conversion Statics.

AN INTRODUCTION TO THE THEORY OF MICROWAVE CIRCUITS

K. KUROKAWA

BELL TELEPHONE LABORATORIES, INC.
MURRAY HILL, NEW JERSEY

6HW 69YFD1160



<17+>0599T66911456400



ACADEMIC PRESS New York and London 1969

large cross-sectional dimensions required. At infrared or optical frequencies, they again become impractical since the cross-sectional dimensions become too small. As a result, the microwave portion of the spectrum is the only region in which ordinary waveguides can be used in practice.

In ordinary waveguides, there are no separate conductors between which the voltage can be defined and the concept of power expressed by the product of voltage and current is no longer directly applicable. Consequently, the electromagnetic field itself has to be investigated first rather than the voltage and current which, in a sense, represent the field. In general, it is hopelessly difficult to solve Maxwell's equations under appropriate boundary conditions and to analyze the electromagnetic field of complex waveguide systems encountered in practice or in laboratories. Even if a formal solution is easily obtainable, the interpretation of the result may become so complicated that no useful information can be extracted from it. In conventional circuit theory, the resistance, capacitance, and inductance are defined without specifying their physical structures or materials used. The relations between the voltage and current at the terminals of such elements are first clarified. Then, the properties of a complex network are studied as a combination of the effects of each element making up the network. There is no need to solve Maxwell's equations explicitly and yet much useful information is obtainable. In a similar fashion, it is possible to study each discontinuity or junction of waveguides separately and to clarify its effect on the propagating modes of the electromagnetic fields either by experiments or by solving Maxwell's equations. The behavior of a complex waveguide system can then be discussed as a combination of the effects of such elements. By this approach, a better understanding can be obtained with less difficulty than trying to solve Maxwell's equations directly for the whole system. The theory of microwave circuits is the study of this approach. The word circuit comes from the similarity existing between this approach and that of conventional circuit theory. The concept of voltage and current can be established by observing the fact that the transverse electric and magnetic fields in a waveguide vary with distance in the same manner as the voltage and current do along a conventional transmission line. Using this concept, impedance can be defined and power appears as the product of voltage and current. Many other useful theorems in conventional ac circuits also become applicable to waveguide systems.

In our category of microwave circuits, we shall include any electromagnetic phenomena in the inside of a hollow region enclosed by conducting walls. We are especially interested in the study of their effect on the propagating

modes of electromagnetic fields in waveguides connected to the hollow region. The actual wavelength is not important; it can be longer or shorter than the centimeter or millimeter ranges as long as the electromagnetic field is confined in a finite region. On the other hand, we shall not discuss waves in free space as in radio and in optics nor those along a dielectric rod or coated wires in open space.

In Chapter 1, some of the topics in the conventional circuit theory are reviewed, which have particular importance in the theory of microwave circuits. These include the theory of transmission lines, bilinear transformations, and power waves. Chapter 2 gives a review of vector analysis and the fundamental properties of electromagnetic fields to facilitate our later study. In Chapter 3, waveguides are discussed for the first time. Here, the eigenvalue problem is studied in detail including the completeness of eigenfunctions. Without the discussion of the completeness, the theory is only partially correct since no answer is given to certain fundamental questions such as why an exponential variation with distance is assumed for each mode and no other possible functional form is considered. Chapter 4 discusses resonant cavities using a similar eigenfunction approach. In Chapter 5, various properties of waveguide junctions are discussed using matrices.

Matrices are introduced since the multitude of symbols and lengthy expressions, which would otherwise be necessary, could interfere with our concentration on the heart of a problem. Just as vectors do, matrices help in studying relatively complex problems, systematically. Chapter 6 discusses two phenomena: One is the coupling between traveling waves and the other is that between electromagnetic fields in cavities. Each illustrates an interesting application of eigenvalue problems. Chapter 7 is a study of linear amplifiers. The central topic is their noise performance. In Chapter 8, we discuss a circuit-theoretical analysis of electron beams which clarifies a number of properties of practical interest. Finally, Chapter 9 contains a discussion of oscillators. Although oscillators are inherently nonlinear, some of their important properties are studied here without specifying the details of the nonlinearity.

Each chapter is followed by a set containing two types of problems: one helps to give deeper understanding of the discussion in the text while the other covers topics which could not be included in the text. The reader is advised to at least read each problem even when he has no time to solve it.

Obviously, this book is not intended to supply microwave circuit design data. Those who want such information should refer to a large collection of books already available on the market (for example, MIT Radiation

Laboratory series). The objective of this book is to convey to the reader, through the discussion of the theory of microwave circuits, some basic ideas and a few mathematical techniques, which are found to be useful in many branches of engineering and science.

In a classroom, since only a few percent of the students will be involved with waveguides and related subjects in the future, there is no reason why details of presently available design techniques should be presented since these will soon be obsolete. On the other hand, basic ideas will ultimately prove to be more useful for solving problems in other fields in which the student may become involved. As a textbook, however, it is not necessary to follow each chapter closely. For example, one could start from Chapter 3 depending on the background of the average student. Furthermore, the teacher does not need to cover the eigenfunction completeness discussion in the classroom; he has only to stimulate the students' interest. They may read the appropriate parts by themselves or refer to them when a similar problem interest them. Similarly, many other parts can be bypassed in the classroom or presented in a simplified form by restricting the discussion to a special case; for example, the reference impedances can be made equal to 1 in the scattering matrices of Chapter 5.

The current volume is based on a first-year graduate text written in Japanese and published by the Maruzen Company in 1963. A number of revisions and additions were made before the book was translated in this final form. The author's cordial thanks are due to the publishing company for allowing him to use material freely for publication.

The English text has been read by many colleagues at Bell Telephone Laboratories. Among them, the author is particularly indebted to Dr. M. R. Barber who accepted the burden of correcting the English. Without his help, this book would not have come into existence.



CONTENTS

<i>Preface</i>	v
--------------------------	---

1. Elements of Conventional Circuit Theory

1.1 Transmission Line Theory.	1
1.2 Smith Chart.	12
1.3 Bilinear Transformations	23
1.4 Power Waves	33
Problems	38

2. Electromagnetic Field Vectors

2.1 Vectors	40
2.2 Maxwell's Equations	63
2.3 Plane Waves	73
Problems	83

3. Waveguides

3.1 Perfect Conductors.	85
3.2 TE ₁₀ Mode in Rectangular Waveguides	87
3.3 Introductory Eigenvalue Problem	99
3.4 Eigenfunctions for Waveguides	113
3.5 General Theory of Waveguides	130
3.6 Examples of Waveguides	138

3.7	Waveguide Discontinuities	149
3.8	Effect of Wall Losses	156
3.9	Waveguides with Inhomogeneous Media	164
	Problems	167
4.	Resonant Cavities	
4.1	Introduction.	170
4.2	Expansions of Electromagnetic Fields	173
4.3	Equivalent Circuits.	185
4.4	Perturbation of Boundaries	193
4.5	Cavities with Inhomogeneous Media	196
	Problems	198
5.	Matrices and Waveguide Junctions	
5.1	Matrices	200
5.2	Reciprocal Conditions	217
5.3	Lossless Conditions	224
5.4	Frequency Characteristic of Lossless Reciprocal Junctions	226
5.5	Lossless Reciprocal Three-port Networks	232
5.6	Symmetrical Y Junctions	234
5.7	Tensor Permeability of Ferrites	240
5.8	Three-port Circulators	247
	Problems	253
6.	Coupled Modes and Periodic Structures	
6.1	Distributed Couplings	257
6.2	Perturbation Method for the Eigenvalue Problem	264
6.3	Interactions between Two Waves.	269
6.4	Periodic Structures	274
6.5	ω - β Diagram	282
	Problems	288
7.	Linear Amplifiers	
7.1	Canonical Forms	290
7.2	Noise Measure	299
7.3	Unconditional Stability	308
7.4	Unilateral Gain.	315
7.5	Tunnel Diode Amplifiers	318

7.6	Manley-Rowe Relations	333
7.7	Negative Resistance Parametric Amplifiers	337
	Problems	343

8. Electron Beams

8.1	An Electron Beam Model	346
8.2	Space Charge Waves	351
8.3	Gap Interactions	359
8.4	Continuous Interaction with a Slow Wave Structure	367
8.5	Electron Beam Noise	372
	Problems	378

9. Oscillators

9.1	A Simplified Oscillator Model	380
9.2	Synchronization by Injection Locking	386
9.3	Noise in Oscillators	389
	Problems	397

Appendix I **Proof of $\lim_{n \rightarrow \infty} k_n^2 = \infty$**

1.	One-Dimensional Case.	399
2.	Two-Dimensional Case	404
3.	Three-Dimensional Case	409
4.	Resonant Cavities with Inhomogeneous Media.	410

Appendix II **Generalized Functions and Fourier Analysis**

References

<i>Subject Index</i>	431
--------------------------------	-----

CHAPTER 1

ELEMENTS OF CONVENTIONAL CIRCUIT THEORY

This chapter reviews several topics in conventional ac circuit theory which are of particular importance in the theory of microwave circuits. They are the theory of transmission lines, Smith chart, bilinear transformations, and power waves. In the discussion of the transmission lines, special emphasis is given to explaining a standard approach in natural science in which an idealized model in one sense or another is constructed and studied in detail. Not only in microwave circuits but also in almost every branch of physics, it is important to construct an appropriate model and to study its behavior in order to obtain a better understanding of the actual situation, which is always far more complicated and difficult to study. Another technique explained in connection with the transmission line theory is that a change in viewpoint of a problem sometimes makes the solution considerably easier. A mathematical operation called linear transformation is one of the methods to change a viewpoint. Since the same kind of operation will be repeatedly employed later in this book, use is explicitly made of a linear transformation, and its meaning and merits are explained. From such a discussion, it will become evident that a problem can become easy or difficult to understand depending on the way it is approached. It is, therefore, of utmost importance for us to keep a flexible attitude and to study various approaches in order to find the simplest techniques.

1.1 Transmission Line Theory

When current i flows through an inductance L_0 , the energy stored in the inductance is magnetic and given by $\frac{1}{2}L_0i^2$. Conversely, if a circuit element

stores magnetic energy when current i is flowing through it, the element has an inductance given by the stored magnetic energy divided by $\frac{1}{2}i^2$. Similarly, if voltage v is applied to a capacitance C_0 , the electric energy stored in the capacitance is given by $\frac{1}{2}C_0v^2$. Therefore, the stored electric energy divided by $\frac{1}{2}v^2$ gives the capacitance.

When direct current i is flowing through a thin wire situated above a conductor ground plane to form a transmission line as shown in Fig. 1.1, there is magnetic field surrounding the wire and, hence, some stored magnetic energy. Except for a small fringing effect at both ends, the magnetic field and, hence, the magnetic energy density is independent of longitudinal position z in Fig. 1.1. Thus, the magnetic energy stored per unit length of

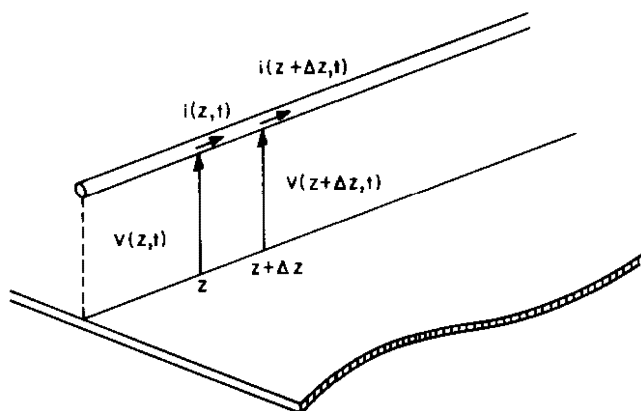


Fig. 1.1. Transmission line above ground plane.

the line is given by the total stored energy divided by the total length of the line. If the energy stored per unit length is further divided by $\frac{1}{2}i^2$, then it should give the inductance L per unit length of the line. On the other hand, if static voltage v is applied between the wire and the ground plane, an electric field appears between them. This means that some electric energy is stored in the space. In a similar way to the above, it is possible to define the electric energy stored per unit length of the line and, hence, the capacitance C per unit length of the line. Thus, L and C defined are constants which are independent of z .

Next, let us consider the case in which a slowly varying voltage is applied to one end of the line instead of static voltage, as before. In this case, one

might expect the voltage and current on the transmission line to be functions of time t as well as of z . Let us first derive the relations between the voltage and current at z and $z + \Delta z$ where Δz is a small increment in z . Since the inductance per unit length of the line is given by L , a length Δz of the line has an inductance $L \Delta z$. The current $i(z, t)$ flowing into this inductance produces potential difference $L \Delta z \{ \partial i(z, t) / \partial t \}$ between z and $z + \Delta z$. The current also changes with z . In this calculation, however, it is assumed that $i(z, t)$ can represent the current at any point between z and $z + \Delta z$ inclusive since Δz is small and later it will be shrunk to zero. The partial differentiation means that the differentiation is with respect to t alone while the current is a function of both z and t . Since the sum of the voltage at $z + \Delta z$ and the potential difference between the terminals of the inductance $L \Delta z$ has to be equal to the voltage at z , we have

$$v(z, t) = v(z + \Delta z, t) + L \Delta z \frac{\partial i(z, t)}{\partial t}$$

or equivalently,

$$- L \Delta z \frac{\partial i(z, t)}{\partial t} = v(z + \Delta z, t) - v(z, t) \quad (1.1)$$

The difference in the currents at z and $z + \Delta z$ is attributable to the current flowing through the shunt capacitance $C \Delta z$ of the length Δz of the line. The shunt current is given by $C \Delta z \{ \partial v(z, t) / \partial t \}$, hence, we have

$$i(z, t) = C \Delta z \frac{\partial v(z, t)}{\partial t} + i(z + \Delta z, t)$$

or equivalently

$$- C \Delta z \frac{\partial v(z, t)}{\partial t} = i(z + \Delta z, t) - i(z, t) \quad (1.2)$$

Dividing (1.1) and (1.2) by Δz and taking the limits of $\Delta z \rightarrow 0$, we have

$$- L \frac{\partial i(z, t)}{\partial t} = \frac{\partial v(z, t)}{\partial z} \quad (1.3)$$

and

$$- C \frac{\partial v(z, t)}{\partial t} = \frac{\partial i(z, t)}{\partial z} \quad (1.4)$$

respectively. By differentiating (1.3) with respect to z and (1.4) with respect to t , $i(z, t)$ can be eliminated. The order of differentiations with respect to t

and z can be reversed without changing the final result. Thus, we obtain a differential equation for $v(z, t)$ alone:

$$LC \frac{\partial^2 v(z, t)}{\partial t^2} = \frac{\partial^2 v(z, t)}{\partial z^2} \quad (1.5)$$

The problem is now reduced to the solution of this equation with appropriate boundary conditions. However, before trying to solve the differential equation, let us consider what assumption have been made. There is an endless list of neglected phenomena, but some of the more important ones are as follows. First, the effect of resistances in the wire as well as in the ground plane is completely neglected. The potential difference between z and $z + \Delta z$ is therefore assumed to be attributable to the inductance $L \Delta z$ alone. Next, if the applied voltage is a function of time, the distribution of current density in the cross section of the wire changes from that for the dc case. The inductance per unit length of the line is therefore different from the dc value, but L is assumed to be a constant in our discussion. Furthermore, when the voltage between conductors becomes sufficiently large, corona discharge and sparks generally take place, and the shunt current can no longer be approximated by $C \Delta z \{\partial v(z, t)/\partial t\}$. Finally, the complicated effect of fringing fields at the ends of the line are ignored.

Any phenomenon that actually takes place is always too complicated to deal with exactly. In order to study it, therefore, a number of things which hopefully produce only minor effects are neglected, and an idealized model is constructed. The phenomenon which takes place in this simplified model is then studied in detail, and the result is compared with the experimentally observed behavior in order to understand actual effects and their causes. Often the first model may not fully explain the actual phenomenon. In such cases, another model has to be constructed which includes one or several factors previously neglected, therefore, it is closer to the actual situation. Sometimes, the same procedure may have to be repeated several times before a satisfactory result is obtained. Selecting an appropriate model constitutes the most important part of any study. If the model is too complicated, a final conclusion may not be obtainable. On the other hand, if it is too simplified, the model may not explain the actual phenomenon at all. Thus, the success or failure of a theory is primarily determined by the selection of an appropriate model.

Now, returning to (1.3)–(1.5), they are the equations which the voltage and current of the idealized model must satisfy. The simplicity of these equations is attributable to the fact that we ignored many, but hopefully

minor, factors in the derivation as explained before. As far as the voltage and current of the model are concerned, we have only to consider these equations. Therefore, let us forget the actual transmission line for the time being, keeping in mind, however, that if the conclusion from these equations fails to explain the actual phenomenon, the actual transmission line must be reconsidered to obtain an improved model.

The differential equation (1.5) is linear: (i) If some voltage v satisfies the differential equation, then a constant times v , Av , also satisfies the same equation; and (ii) If v_1 and v_2 satisfy the differential equation separately, then their sum $v_1 + v_2$ also satisfies the equation. It follows from (i) that if the voltage applied to the model transmission line is increased by a certain factor, the voltage at every point along the line will be increased by the same factor. (ii) asserts that the principle of superposition holds for the model. For instance, when the applied voltage has a complex waveform, it can be decomposed into simpler components, and the line voltage can be calculated for each separately. The line voltage corresponding to the original complicated applied voltage can then be obtained by summing the results thus obtained. Another important point in connection with (1.5) is that the coefficient of the differential equation is real, i.e., the constant LC is not a complex quantity but rather a real quantity.

Since Eq. (1.5) is a linear differential equation with real coefficients, we can use a special function of time, $e^{j\omega t}$. Suppose $v_1 = V(z) e^{j\omega t}$ satisfies the differential equation, then the complex conjugate $v_2 = V^*(z) e^{-j\omega t}$ is also seen to satisfy the same differential equation since LC is real. Because of the linearity, a constant times their sum as well as their difference will also satisfy (1.5):

$$\begin{aligned}\sqrt{2} \{V(z) e^{j\omega t} + V^*(z) e^{-j\omega t}\} / 2 &= \text{Re} \{ \sqrt{2} V(z) e^{j\omega t} \} \\ -j \sqrt{2} \{V(z) e^{j\omega t} - V^*(z) e^{-j\omega t}\} / 2 &= \text{Im} \{ \sqrt{2} V(z) e^{j\omega t} \}\end{aligned}$$

where the asterisk (*) indicates the complex conjugate, and Re and Im indicate the real and imaginary parts of the quantity following. The quantity $\sqrt{2}$ is introduced in front of $V(z)$ so that $V(z)$ represents the effective value as is familiar in the ac circuit theory. The above two quantities are both real functions of time, and they can be interpreted as the actual voltages which are measurable on the model transmission line. Therefore, let us agree to take either the real or the imaginary part of $\sqrt{2} e^{j\omega t}$ times $V(z)$ for the interpretation of $V(z)$. Since ω is arbitrary, by the theory of Fourier integral, the result can be integrated with respect to ω with a proper weighting function to obtain the line voltage corresponding to any time-varying

applied voltage. For this reason, we shall use $e^{j\omega t}$ as the time factor unless otherwise specified. With the time factor $e^{j\omega t}$, the differentiation with respect to time becomes equivalent to the multiplication by $j\omega$ thereby simplifying the calculation. If $\sin \omega t$ is adopted as the time factor, the differentiation will produce a term like $\omega \times \cos \omega t$ and the calculation will not be as simple as in the case in which $e^{j\omega t}$ is used.

Substituting $v = V(z)e^{j\omega t}$ into (1.5) and dividing both sides by $e^{j\omega t}$, we obtain the differential equation for $V(z)$:

$$-\omega^2 LCV(z) = \frac{d^2 V(z)}{dz^2} \quad (1.6)$$

Since $V(z)$ is not a function of t , the partial derivative with respect to z is changed to an ordinary derivative. Since (1.6) is a linear ordinary differential equation, it has a solution of the form

$$V(z) = Ae^{-\gamma z} \quad (1.7)$$

where A and γ are constants. The substitution of (1.7) into (1.6) gives the condition for (1.7) to be a solution,

$$\begin{aligned} \gamma^2 &= -\omega^2 LC \\ \gamma &= \pm j\omega(LC)^{1/2} = \pm j\beta \end{aligned} \quad (1.8)$$

Here, in order to eliminate possible confusion about the sign of square roots, let us agree that $()^{1/2}$ expresses a positive value. The negative root is expressed by $-()^{1/2}$. The constant γ is called the propagation constant, and β the phase constant.

From (1.7) and (1.8), it follows that

$$V(z) = Ae^{-j\beta z} + Be^{j\beta z} \quad (1.9)$$

is also a solution of (1.6). A and B are constants to be determined by boundary conditions. Since (1.6) is a second order ordinary differential equation and (1.9) has two constants to be determined, this is the most general form of the solution and no other functions have to be examined. The proof of this assertion for our special case is not difficult. Let $V(z)$ be any solution of (1.6) with nonzero $\gamma^2 = -\omega^2 LC$. We shall assume that A and B are not constant but that they are functions of z satisfying the following equations.

$$V(z) = A(z)e^{-\gamma z} + B(z)e^{\gamma z} \quad (1.10)$$

$$V'(z) = A(z)(-\gamma e^{-\gamma z}) + B(z)\gamma e^{\gamma z} \quad (1.11)$$

where the prime indicates the derivative with respect to z . Since the determinant of the coefficients of $A(z)$ and $B(z)$ is not equal to zero, it is always possible to solve for $A(z)$ and $B(z)$ for a given $V(z)$. From (1.10), $V'(z)$ is calculated as

$$V'(z) = A(z)(-\gamma e^{-\gamma z}) + B(z)\gamma e^{\gamma z} + A'(z)e^{-\gamma z} + B'(z)e^{\gamma z}$$

Since the first two terms on the right-hand side represent $V'(z)$ as given by (1.11), we have

$$A'(z)e^{-\gamma z} + B'(z)e^{\gamma z} = 0 \quad (1.12)$$

From (1.11), $V''(z)$ can be calculated:

$$V''(z) = A(z)\gamma^2 e^{-\gamma z} + B(z)\gamma^2 e^{\gamma z} + A'(z)(-\gamma e^{-\gamma z}) + B'(z)\gamma e^{\gamma z}$$

The first two terms on the right-hand side represent $\gamma^2 V(z)$. However, since $V(z)$ is a solution of the differential equation $V''(z) = \gamma^2 V(z)$, we must have

$$A'(z)(-\gamma e^{-\gamma z}) + B'(z)\gamma e^{\gamma z} = 0 \quad (1.13)$$

From (1.12) and (1.13), both $A'(z)$ and $B'(z)$ are found to be zero. Thus, we conclude that any solution of the differential equation can be expressed in the form (1.9) with A and B being constant.

From (1.4), it follows that $i(z, t)$ can also have the time factor $e^{j\omega t}$. Therefore, let us write $i(z, t) = I(z)e^{j\omega t}$. Substituting this and (1.9) into (1.3) we have

$$-j\omega LI(z) = -j\beta Ae^{-j\beta z} + j\beta Be^{j\beta z}$$

Using the relation $\beta = \omega(LC)^{1/2}$, the above equation can be rewritten in the form

$$I(z) = Z_0^{-1}(Ae^{-j\beta z} - Be^{j\beta z}) \quad (1.14)$$

where

$$Z_0 = (L/C)^{1/2} = Y_0^{-1} \quad (1.15)$$

Since $\omega L \Delta z$ has the dimension of impedance and $\omega C \Delta z$ that of admittance (where L = inductance per unit length), $(L/C)^{1/2}$ has the dimension of impedance. It is, therefore, indicated by Z_0 and called the characteristic impedance of the line, its inverse Y_0 is called the characteristic admittance.

It follows from (1.9) and (1.14) that the voltage and current on the line are completely determined when the two constants A and B are fixed. This means that the voltage and current at every point on the line are automatically determined if the following conditions are specified:

- (i) both voltage and current are given at one point along the line;

- (ii) the voltages are given at two different points a distance l apart where $\beta l \neq n\pi$ (n : integer);
- (iii) the currents at these two points are given; and
- (iv) the ratio of the voltage to the current at one point and either the voltage or current at another are specified.

Since both A and B are complex numbers, four real numbers are specified in the above procedure.

Up to this point, our discussion has been based on the voltage and current. Let us now introduce the following quantities

$$a(z) = \frac{1}{2}Z_0^{-1/2} \{V(z) + Z_0 I(z)\}, \quad b(z) = \frac{1}{2}Z_0^{-1/2} \{V(z) - Z_0 I(z)\} \quad (1.16)$$

which are the result of a linear transformation applied to $V(z)$ and $I(z)$. With Z_0 being fixed, both $a(z)$ and $b(z)$ are determined if $V(z)$ and $I(z)$ are given. Conversely, if $a(z)$ and $b(z)$ are given, from the inverse transformation

$$V(z) = Z_0^{1/2} \{a(z) + b(z)\}, \quad I(z) = Z_0^{-1/2} \{a(z) - b(z)\} \quad (1.17)$$

both $V(z)$ and $I(z)$ can be obtained. Consequently, there should be no difference in the result whether the transmission line is studied in terms of $V(z)$ and $I(z)$ or in terms of $a(z)$ and $b(z)$. However, since $a(z)$ and $b(z)$ can be shown to have physical meaning which helps in the understanding of transmission line phenomena, it is worthwhile to introduce them. Before

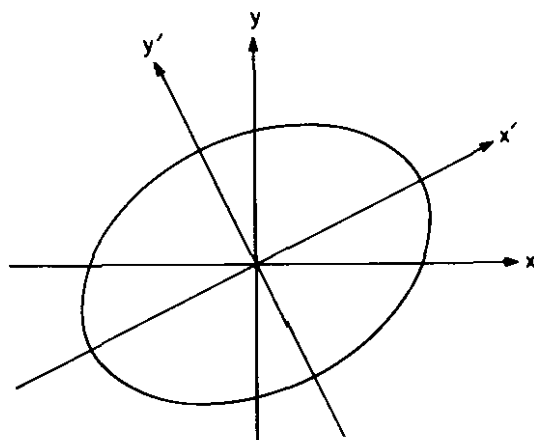


Fig. 1.2. An example of coordinate-transformations.

explaining the physical meaning, let us look at a familiar example in which a linear transformation facilitates our study. Referring to Fig. 1.2, various properties of an ellipse can be studied in terms of the coordinates (x, y) as well as in terms of (x', y') which are obtained by a linear transformation of (x, y) . Since everything that can be explained in terms of (x', y') can be explained in terms of (x, y) , the linear transformation is not entirely necessary in the study of the ellipse. However, since the coordinates (x', y') give a simpler expression for the ellipse, properties which are not clear in terms of (x, y) may be disclosed in terms of (x', y') . For this reason (x', y') has an advantage over (x, y) , and one should be familiar with this kind of transformation.

Let us now substitute (1.9) and (1.14) into (1.16) and investigate the physical meaning of $a(z)$ and $b(z)$. The results are

$$a(z) = Z_0^{-1/2} A e^{-j\beta z}, \quad b(z) = Z_0^{-1/2} B e^{j\beta z} \quad (1.18)$$

$a(z)$ and $b(z)$ are directly related to A and B , respectively. It is worth noting that the magnitudes of both $a(z)$ and $b(z)$ are kept constant along the line. Their phases change but the angles are directly proportional to z . Thus, the variations of $a(z)$ and $b(z)$ with z are considerably simpler than those of $V(z)$ and $I(z)$.

In the lower frequency ranges, a piece of wire is generally used to connect circuit elements together, and in this wire the voltage and current are kept constant. Voltage and current are primarily used for the study of networks in which many elements are connected together by means of conducting wires because of this simple relation of constancy. In the higher frequency ranges, the connections are made through transmission lines. When their length can no longer be neglected, the voltage and current vary with z in the complicated manner shown in (1.9) and (1.14). It is natural, therefore, to look for substitutes which are the most simple functions of z as possible while providing the same amount of information as $V(z)$ and $I(z)$. In order to satisfy the last condition, quantities obtainable through a linear transformation of $V(z)$ and $I(z)$ are considered, and $a(z)$ and $b(z)$, which have constant magnitudes and phases directly proportional to z , are chosen. Because of these simple relations, $a(z)$ and $b(z)$ will be found particularly suitable for the study of transmission line phenomena.

For the interpretation of $a(z)$ when quantities are varying sinusoidally with time, we have to take the real or imaginary part of $\sqrt{2} e^{j\omega t}$ times $a(z)$ just as we did for the interpretation of $V(z)$. Let us consider the real part.

Referring to Fig. 1.3, we obtain using (1.18)

$$\begin{aligned} \operatorname{Re} \{ \sqrt{2} a(z) e^{j\omega t} \} &= \sqrt{2} Z_0^{-1/2} \operatorname{Re} \{ A e^{j(\omega t - \beta z)} \} \\ &= \sqrt{2} Z_0^{-1/2} |A| \cos(\omega t - \beta z + \varphi) \end{aligned} \quad (1.19)$$

where

$$\tan \varphi = \operatorname{Im} \{ A \} / \operatorname{Re} \{ A \} \quad (1.20)$$

With fixed t , the right-hand side of (1.19) represents a sinusoidal waveform as shown in Fig. 1.4. The wavelength λ is equal to $2\pi/\beta$. If $\Delta z = (\omega/\beta) \Delta t$,

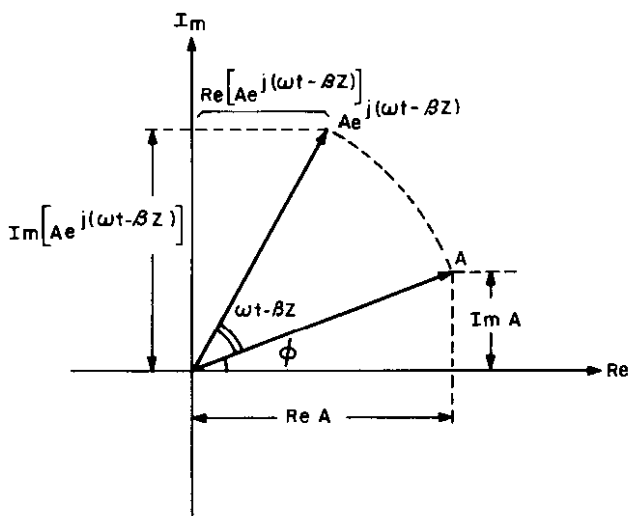


Fig. 1.3. Relation between two complex vectors A and $Ae^{j(\omega t - \beta z)}$.

the arguments of the cosine are equal at (z, t) and $(z + \Delta z, t + \Delta t)$, as we can easily see from

$$\begin{aligned} \omega t - \beta z + \varphi &= \omega(t + \Delta t) - \beta \{ z + (\omega/\beta) \Delta t \} + \varphi \\ &= \omega(t + \Delta t) - \beta(z + \Delta z) + \varphi \end{aligned}$$

Since z is arbitrary, this means that after the time interval Δt , the function as a whole is translated toward the positive z -direction by the amount Δz as shown by the dotted line in Fig. 1.4. The velocity of the translation is given by

$$v_p = (\Delta z / \Delta t) = (\omega / \beta) \quad (1.21)$$

From these observations, we conclude that (1.19) represents a wave moving

toward the positive z -direction with a constant velocity v_p which is called the phase velocity. In our case, since $\beta = \omega(LC)^{1/2}$,

$$v_p = \omega/\beta = (LC)^{-1/2} \quad (1.22)$$

Thus, we see that $a(z)$ is a wave moving in the positive z -direction with the phase velocity $(LC)^{-1/2}$. Similarly, $b(z)$ represents a wave moving in the

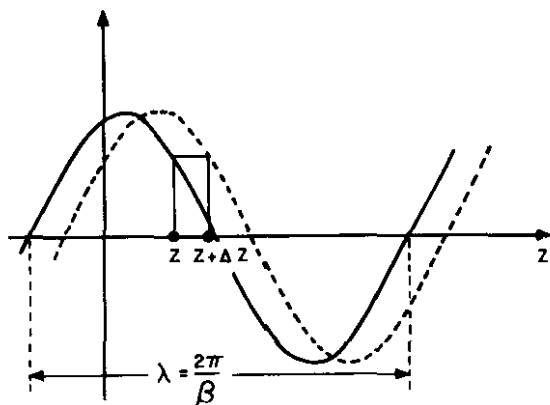


Fig. 1.4. Explanation for the wave motion of $a(z)$.

negative z -direction with the same phase velocity. For this reason, they are called traveling waves.

Finally, let us calculate the net power at z flowing toward the positive z direction. Referring to Fig. 1.5, the power is given by

$$\begin{aligned} P &= \text{Re} \{ V(z) I^*(z) \} = \text{Re} [\{ a(z) + b(z) \} \{ a^*(z) - b^*(z) \}] \\ &= \text{Re} \{ |a(z)|^2 - |b(z)|^2 + a^*(z) b(z) - a(z) b^*(z) \} \end{aligned}$$

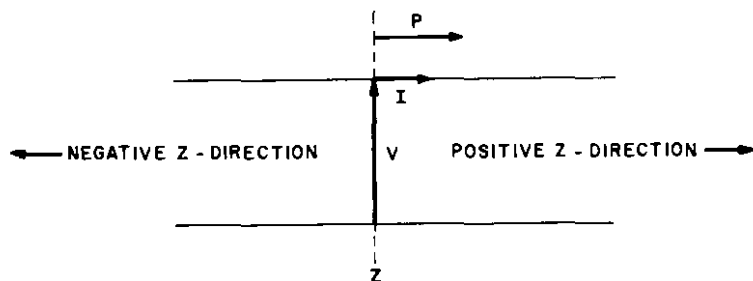


Fig. 1.5. Relation between V , I , and P .

where use is made of (1.17). Since $a^*(z)b(z) - a(z)b^*(z)$ is the difference between a complex number and its conjugate, it is always imaginary. Therefore, we have

$$P = |a(z)|^2 - |b(z)|^2 \quad (1.23)$$

This result can be interpreted as follows. The wave $a(z)$ moving in the positive z -direction carries the power $|a(z)|^2$, and similarly the wave $b(z)$ carries the power $|b(z)|^2$ toward the negative z -direction. The net power toward the positive z -direction is therefore given by $|a(z)|^2 - |b(z)|^2$. In other words, each wave defined above is accompanied by the power equal to the square of its magnitude.

The advantages of introducing $a(z)$ and $b(z)$ will be further exploited in connection with the following discussion of the Smith chart.

1.2 Smith Chart

When we study electric circuits in terms of voltage and current, their ratio, i.e., impedance, proves to be a useful quantity. Since voltage and current have been replaced by $a(z)$ and $b(z)$, let us take the ratio of the latter quantities:

$$r = b(z)/a(z) = (B/A) e^{j2\beta z} \quad (1.24)$$

This is called the reflection coefficient. At a reference plane z , if $a(z)$ is considered to be an incident wave, $b(z)$ represents the reflected wave, and r expresses the magnitude and phase of the reflected wave relative to the incident wave. In other words, r is the reflection corresponding to a unit-incident wave. When the reference plane z is shifted toward the positive z -direction with a constant velocity, r rotates counterclockwise in the complex plane with an uniform speed without changing its magnitude as we can see from (1.24). When z changes by $\pi/\beta = \lambda/2$, r completes one rotation. If there is no power source on the right-hand side of the reference plane z , the net power flow toward the positive z -direction must be positive and hence $|a(z)|^2 \geq |b(z)|^2$ from (1.23). This means that $|r|$ is smaller than unity in cases in which the transmission line is terminated by a passive circuit.

Although in Eq. (1.9), the variation of the magnitude of the line voltage with z appears, at first, to be complicated, the use of r provides a considerable simplification. From (1.17), (1.18), and (1.24), we have

$$|V(z)| = |Z_0^{1/2} \{a(z) + b(z)\}| = Z_0^{1/2} |a(z)| |1 + r| = |A| |1 + r| \quad (1.25)$$

$|V(z)|$ is, therefore, proportional to $|1 + r|$. Since 1 is a constant and r

uniformly rotates with z , the vector diagram shown in Fig. 1.6 clarifies how $|V(z)|$ varies with z . The reflection coefficient r rotates once as z shifts by $\lambda/2$ hence the variation of $|V(z)|$ repeats itself every half wavelength as shown by the solid line in Fig. 1.7. Similarly, from (1.17), the magnitude of the line current is given by

$$\begin{aligned} |I(z)| &= |Z_0^{-1/2} \{a(z) - b(z)\}| = Z_0^{-1/2} |a(z)| |1 - r| \\ &= Z_0^{-1} |A| |1 - r| \end{aligned} \quad (1.26)$$

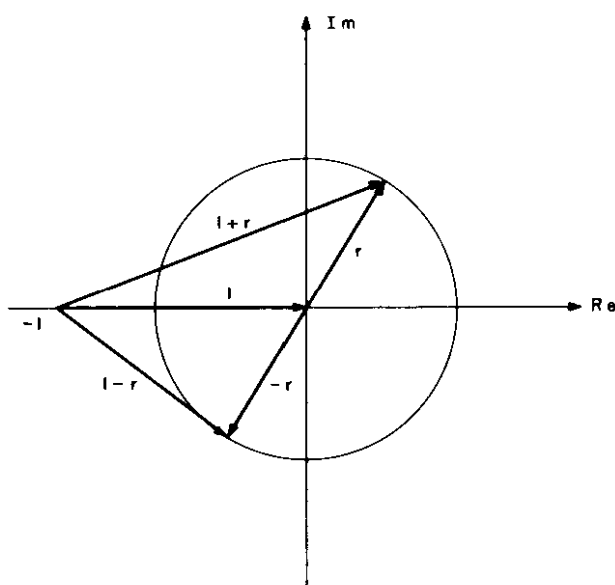


Fig. 1.6. Vectors $1 + r$ and $1 - r$.

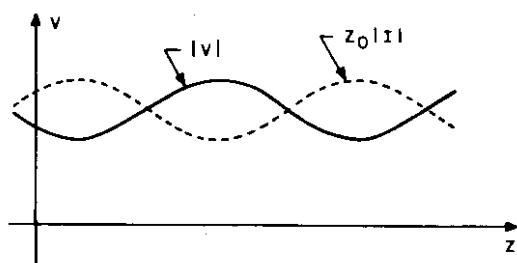


Fig. 1.7. Relative amplitudes of voltage and current.

Referring to Fig. 1.6, it is clear that the magnitude of current changes with z as shown by the dotted line in Fig. 1.7. At the maximum point of $|V(z)|$, $|I(z)|$ becomes minimum and vice versa. Except for the fact that they are shifted by $\lambda/4$ relative to each other, the two curves have identical shape.

As we have explained in connection with (1.16) and (1.17), effects which can be discussed in terms of voltage and current can be explained in terms of $a(z)$ and $b(z)$. Consequently, concept of impedance may seem redundant once the reflection coefficient is introduced. Inasmuch as we are already familiar with the concept of impedance at low frequencies and because the series or parallel connection of circuit elements can be expressed by the simple addition of their impedances or admittances, respectively, it is advantageous to use freely both the reflection coefficient and impedance, switching from one quantity to another during a course of circuit study. For this purpose, let us clarify the relation between the reflection coefficient r and impedance Z . By definition, the impedance is given as the ratio of $V(z)$ to $I(z)$. Therefore using (1.17), we have

$$Z = \frac{V(z)}{I(z)} = \frac{Z_0^{1/2} \{a(z) + b(z)\}}{Z_0^{-1/2} \{a(z) - b(z)\}} = Z_0 \frac{1+r}{1-r}$$

or equivalently,

$$Z/Z_0 = (1+r)/(1-r) \quad (1.27)$$

The left-hand side of (1.27) is a dimensionless quantity representing the impedance Z relative to Z_0 . This is called the normalized impedance. Solving (1.27) for r , the reflection coefficient can be expressed in terms of the normalized impedance:

$$r = \frac{(Z/Z_0) - 1}{(Z/Z_0) + 1} \quad (1.28)$$

From (1.27) and (1.28), we see that there is a one-to-one correspondence between r and Z/Z_0 . If Z/Z_0 is plotted on an r plane or if r is plotted on the Z/Z_0 plane, when one is given, the other can be obtained conveniently from the chart. When the reference plane is shifted, a new reflection coefficient can be obtained on the r plane simply by rotating the original vector r by the angle corresponding to the distance between the new and old reference planes. Such a simple operation is not available for obtaining the new Z/Z_0 on the Z/Z_0 plane. For this reason, a plot of Z/Z_0 on the r plane is preferred. Among the various ways to plot Z/Z_0 , the most common is to draw two sets of curves: constant resistance and constant reactance. The real and imaginary parts of Z/Z_0 corresponding to a given r can be read

on the chart from the two curves intersecting at the point representing r . The reflection coefficient plane on which the normalized impedance is plotted in this way is called a Smith chart.

In order to construct the Smith chart, let us first consider the constant resistance locus on the r plane, i.e., the image on the r plane of a straight line parallel to the imaginary axis on the Z/Z_0 plane, as shown by K_0 in Fig. 1.8(a). The strategy, here, is to decompose the transformation (1.28)

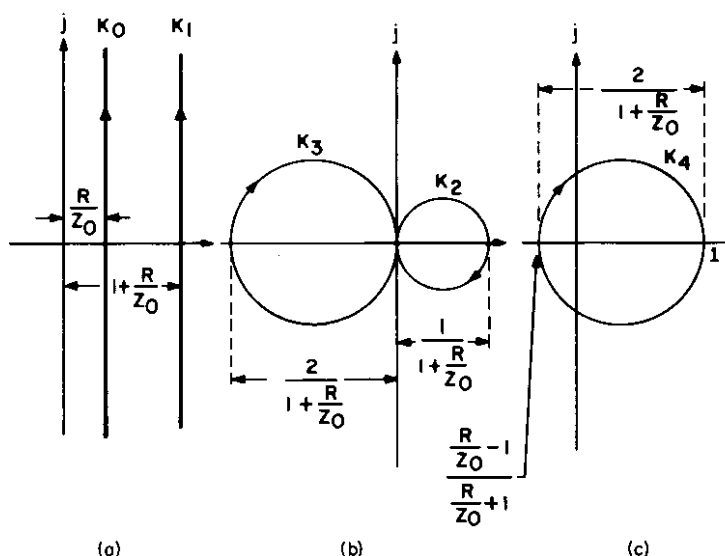


Fig. 1.8. Drawing of Smith chart (1).

into several simpler ones and to study the steps one by one. Equation (1.28) can be rewritten in the form

$$r = 1 - \frac{2}{(Z/Z_0) + 1} \quad (1.29)$$

The first step is to determine $\{(Z/Z_0) + 1\}$. For this transformation, each point on K_0 is shifted to the right by a unit distance giving K_1 in Fig. 1.8(a). We now use the result that the inverse of a complex number is expressed on

the complex plane by a vector with an inverse magnitude and having an angle which is the negative of the original one. The inverse of each point on K_1 can therefore be plotted, resulting in circle K_2 shown in Fig. 1.8(b). This is $1/\{(Z/Z_0) + 1\}$. The proof that the inverse of a straight line is a circle is given in Section 1.3. Multiplying K_2 by -2 , we obtain circle K_3 which corresponds to the second term on the right-hand side of (1.29). By adding 1 to K_3 , i.e., shifting K_3 to the right by a unit distance, the locus on the r plane is obtained which corresponds to a constant resistance line on the Z/Z_0 plane. This is circle K_4 in Fig. 1.8(c). For different values of R/Z_0 , the radius of the circle changes. However, regardless of the value of R/Z_0 , the circle is symmetrical with respect to the real axis and always passes through the point $(1, j0)$.

Next, let us consider the constant reactance locus. The straight line C_0 in Fig. 1.9(a) for which X/Z_0 is constant coincides with itself when trans-

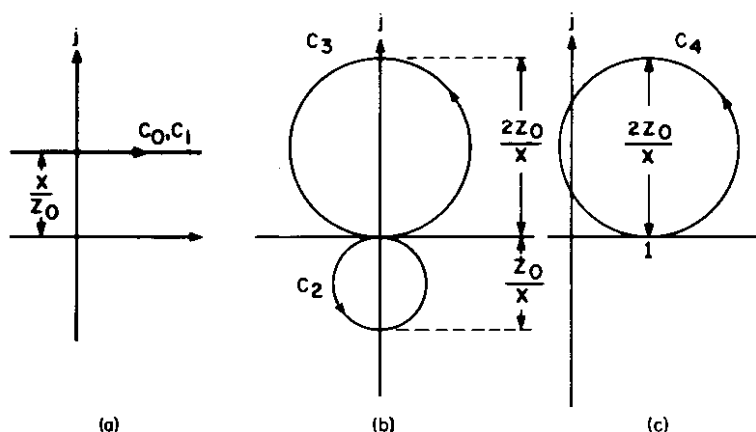


Fig. 1.9. Drawing of Smith chart (2).

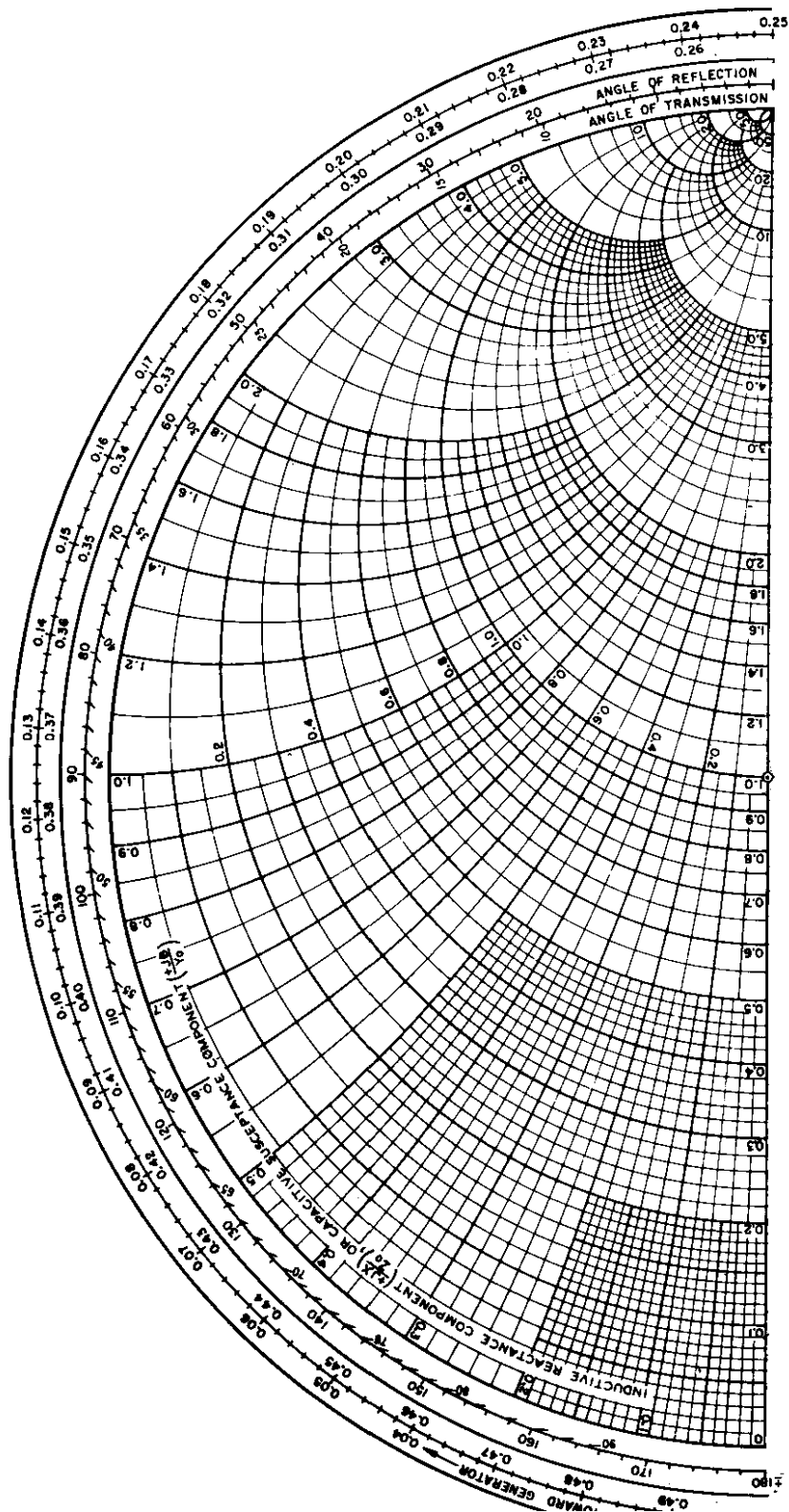
formed to $\{(Z/Z_0) + 1\}$ as indicated by C_1 . The inverse of C_1 is circle C_2 , shown in Fig. 1.9(b), which represents $1/\{(Z/Z_0) + 1\}$, and from this $-2/\{(Z/Z_0) + 1\}$ is given by C_3 in Fig. 1.9(b). Shifting C_3 to the right by a unit distance, the desired locus on the r plane is obtained as shown by C_4 in Fig. 1.9(c). For different values of X/Z_0 , the radius of the circle varies. For a negative value of X/Z_0 , the circle appears below the real axis. However, regardless of the value of X/Z_0 , the circle passes through the point $(1, j0)$ and is tangent to the real axis.

Thus, two sets of circles corresponding to various values of R/Z_0 and X/R_0 can be drawn on the r plane. Excluding the region in which $|r| > 1$, we obtain the Smith chart shown in Fig. 1.10. The region where $|r| > 1$ is excluded because the transmission line is ordinarily terminated by a passive network for which $|r|$ is less than unity, as we explained before. It follows from Fig. 1.8 that the excluded part in the r plane of Fig. 1.8(c) corresponds to the left-hand side of the imaginary axis of the Z/Z_0 plane which is shown in Fig. 1.8(a), i.e., to the region where the resistance is negative. In addition to these two sets of curves, the Smith chart, shown in Fig. 1.10, has scales around it showing angles in the corresponding distances in wavelengths. These are convenient for obtaining the reflection coefficients at different points along the transmission line when the distance is measured in wavelengths.

To illustrate how to use the Smith chart, let us discuss the case in which a load impedance Z_L is connected at the far end of a transmission line. First Z_L/Z_0 is calculated and the point P corresponding to Z_L/Z_0 is located on the Smith chart. The vector OP with its tip at P and its tail at the origin O on the r plane gives the reflection coefficient r_L at the reference plane where Z_L is connected. Let OP in Fig. 1.11 be the vector r_L . If the reference plane is shifted by k wavelengths toward the generator, i.e., toward the negative z -direction, the reflection coefficient vector r rotates clockwise by k on the scale around the Smith chart. The reflection coefficient at the new reference plane is given by OQ in Fig. 1.11, and Z/Z_0 corresponding to Q gives the normalized impedance looking into the load at the new reference plane. The actual impedance is then obtained by multiplying this value by Z_0 . As the reference plane shifts further toward the generator, r rotates, and at a certain point the vector r lies on the real axis. This means that the impedance looking into the load from this point becomes purely resistive. Further movement of the plane by $\lambda/4$ gives another pure resistance with the normalized value being the inverse of the previous one. The magnitude of the line voltage is proportional to $|1 + r|$ which is equal to the length of the vector drawn from the point $r = -1 + j0$ to the tip of vector r . Note that $r = -1 + j0$ corresponds to $Z/Z_0 = 0 + j0$. In the case of Fig. 1.11, as the reference plane moves toward the generator from the load, the voltage first increases. It takes the maximum value when r lies on the real axis and then decreases to the minimum point located a quarter wavelength from the maximum. A further shift increases the voltage again and thereafter repeats the cycle.

The ratio of the maximum to the minimum voltage is called the standing

IMPEDANCE OR ADMITTANCE COORDINATES



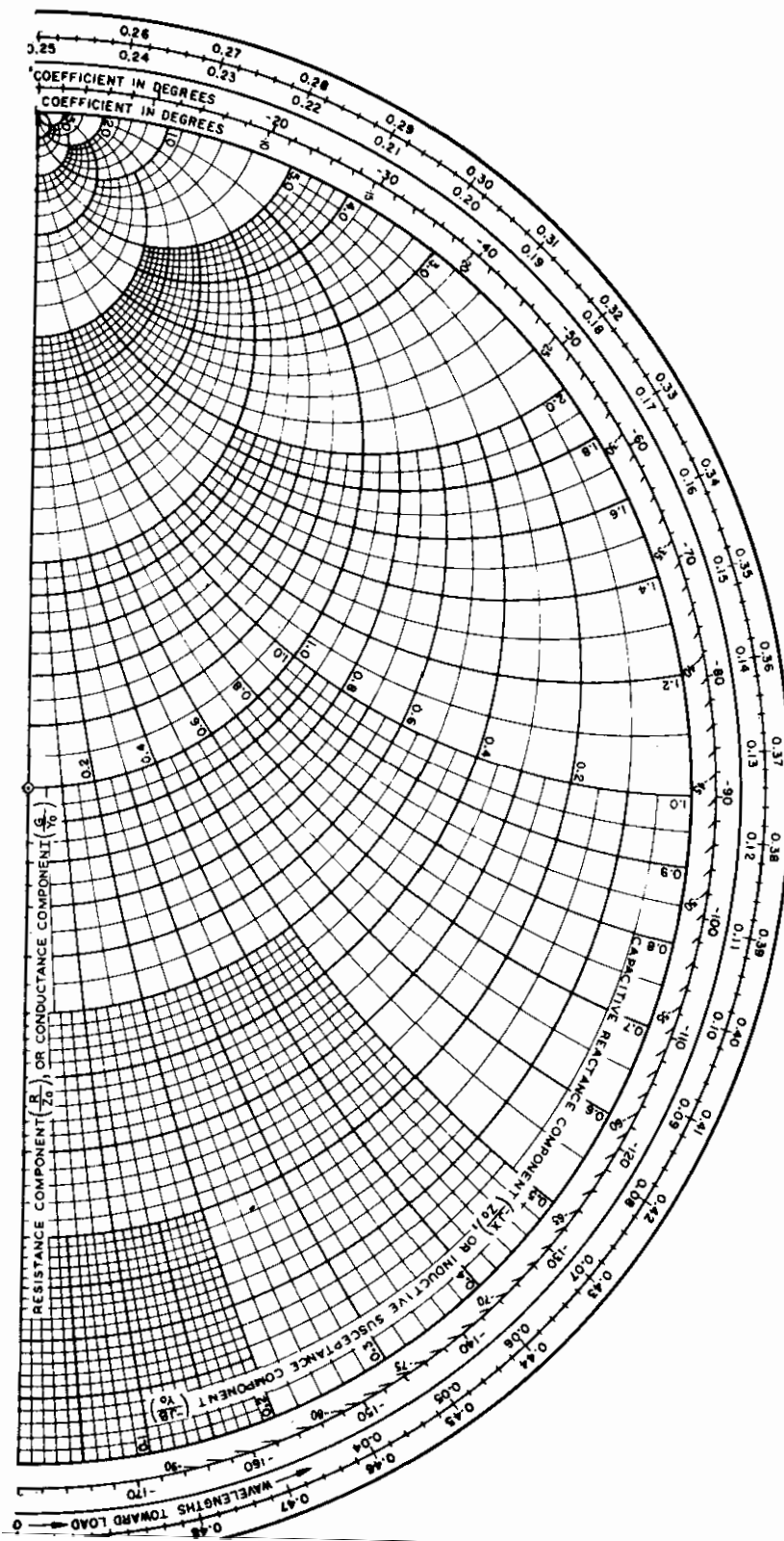


Fig. 1.10. Smith chart. (Reproduced by permission from Kay Electric Company, Pine Brook, New Jersey.)

wave ratio (SWR). From Fig. 1.11, the standing wave ratio is given by

$$\text{SWR} = \frac{1 + |r_L|}{1 - |r_L|} \quad (1.30)$$

where r_L is the reflection coefficient of the load. A comparison of (1.27) with (1.30) shows that the standing wave ratio is given by the normalized impedance at the point where r is equal to $|r_L|$. To obtain the value of standing wave ratio, therefore, draw a circle with center at the origin O and passing through the point P . It crosses the real axis at two points. Take the right-hand point and read the normalized resistance. This gives the desired standing wave ratio.

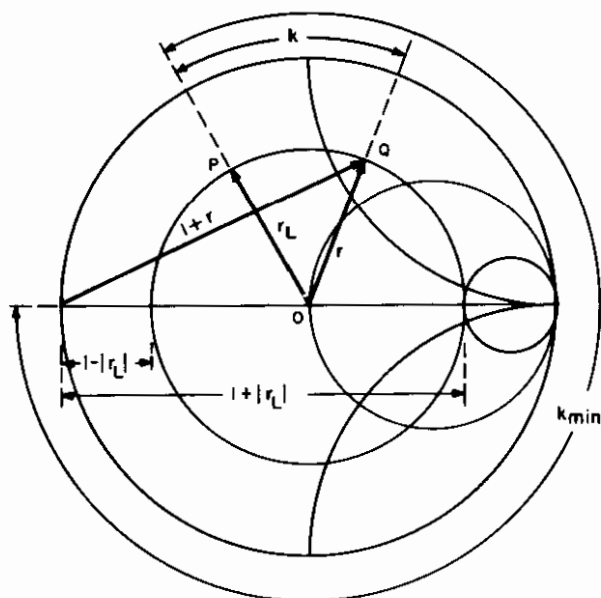


Fig. 1.11. Example of using Smith chart.

Now suppose that the standing wave ratio and the distance from the load to the first voltage minimum point are given and the load impedance is to be calculated. To obtain the normalized impedance, first draw a circle with center at the origin passing through the point whose normalized impedance is equal to the standing wave ratio. Next, measure the given distance counterclockwise from the $(0 + j0)$ impedance point on the scale around the

Smith chart and draw a straight line from there to the origin. This follows since minimum voltages always occur at points whose position corresponds to $Z/Z_0 = 0 + j0$. The intersection of this line and the circle previously drawn gives the reflection coefficient and, hence, the normalized impedance of the load. In this way, the normalized impedance can be obtained once the SWR along a transmission line and the voltage minimum point have been measured. This is called the standing wave method and is one of the most commonly used methods of measuring high frequency impedances. Once the characteristic impedance is calculated or calibrated against some standard impedance, the unknown impedance can be determined.

So far, the voltage and impedance have been almost exclusively used. Essentially the same procedure can be developed using current and admittance, but this will be omitted except for some comments concerning the following useful relations. The inverse of the normalized impedance is called the normalized admittance, since

$$(Z/Z_0)^{-1} = Z_0/Z = Y/Y_0$$

Taking the inverse of (1.27), we have

$$Y/Y_0 = (1 - r)/(1 + r) \quad (1.31)$$

A comparison of (1.27) with (1.31) shows that if r is replaced by $-r$, the normalized impedance is transformed into the normalized admittance. Since the normalized impedance is given at the point r on the Smith chart, the normalized admittance is obtained at the point $-r$. In other words, given a normalized impedance on the Smith chart, the normalized admittance is obtained at the opposite point with respect to the origin. This means that the same chart can be used for admittance as well as for impedance calculations. The point $(-1, j0)$ on the r plane corresponds to the normalized impedance $0 + j0$ and also to the normalized admittance $\infty + j\infty$.

At this point, in order to become more familiar with the Smith chart, let us solve a particular problem.

Exercise: Suppose that a $50\ \Omega$ transmission line is terminated by a $30\ \Omega$ resistor. By connecting a capacitor in shunt at an appropriate point along the line, eliminate the reflection toward the generator.

Answer: The normalized impedance of the load is given by

$$Z/Z_0 = (30 + j0)/50 = 0.6 + j0$$

Let us draw a circle with center at the origin passing through the point

whose normalized impedance is $0.6 + j0$. The normalized impedance looking into the load at any point along the line should be found on this circle. The normalized admittance is also located on the same circle. We wish to eliminate the reflection by inserting a capacitor in shunt. To do this, on the same circle, we must find a point whose normalized admittance plus jb ($b > 0$) is equal to $1 + j0$, implying no reflection where b is the normalized susceptance of the capacitor to be determined. The desired point corresponds to the normalized admittance $1 - jb$. Therefore, the point we are seeking is one of the intersections between the above circle and the unit conductance circle. Since $b > 0$, the imaginary part must be negative. Thus, the normalized admittance is determined on the Smith chart as $1 - j0.5$. The opposite

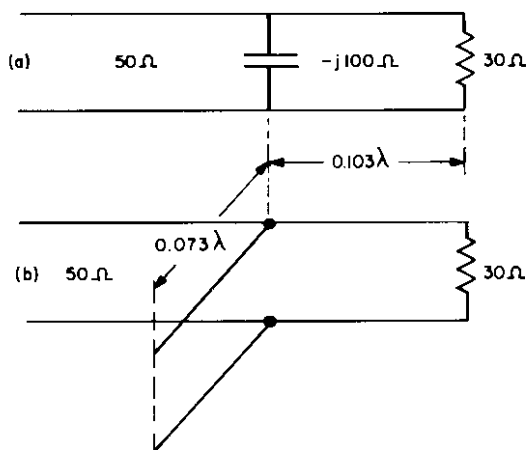


Fig. 1.12. Example of matching circuit.

point with respect to the origin gives the normalized impedance. On the scale around the Smith chart, this point is found to be about 0.103 wavelength away from the load. In other words, if a capacitor, whose normalized admittance is $j0.5$, is inserted in shunt at a point 0.103 wavelength from the resistor, the reflection is canceled out and the generator sees a matched load. The normalized impedance of the capacitor is $-j2$ and the impedance $-j100\ \Omega$. The circuit configuration is shown in Fig. 1.12(a). The same normalized admittance $j0.5$ can also be obtained using an open-ended $50\ \Omega$ transmission line about 0.073 wavelength long which is easily seen using the Smith chart. Thus, Figs. 1.12(a) and (b) are equivalent to each other.

In the above discussion, the losses of the transmission line have been

completely neglected. However, it should be remembered that because of the resistance in the wire and in the ground plane and because of the possible dielectric loss in the medium, the magnitudes of the waves along an actual transmission line decrease as they progress. If the losses are taken into account, the propagation constant γ which was defined in (1.7) becomes a complex quantity $\alpha + j\beta$ and, at the same time, a nonzero imaginary part appears in Z_0 . The real part of γ , α , is called the attenuation constant. When the losses are small, both α and the imaginary part of Z_0 are small and of the same order of magnitude. The effect of α , however, appears in the exponential form $e^{-\alpha z}$, and it becomes larger as $|z|$ increases; in other words, α is a measure of the loss per unit length. On the other hand, the effect of the imaginary part of Z_0 has no such enhancement with increasing $|z|$. For these reasons, Z_0 is generally assumed to be real, and γ complex. In this case,

$$r = (B/A) e^{2\alpha z + j2\beta z} \quad (1.32)$$

The locus of r , therefore, becomes a spiral instead of a circle, with a radius which increases with increasing z . If the reference plane is moved away from the load, $|r|$ gets smaller. When the generator is sufficiently far away from the load, it sees no reflection regardless of the load impedance at the other end. If r is given at z_1 , and r is required at another point z_2 , the magnitude has to be multiplied by $e^{2\alpha(z_2 - z_1)}$ while the phase must be rotated by $e^{j2\beta(z_2 - z_1)}$. Provided that r is transformed in this manner, the normalized impedance as well as the admittance can be obtained in the same manner as in the previous lossless case. However, the magnitude of the voltage or current is no longer proportional to $|1 + r|$ or $|1 - r|$, respectively, as z varies.

1.3 Bilinear Transformations

A complex variable w is described as being a bilinear transformation of another complex variable z when w is expressed in the form

$$w = \frac{az + b}{cz + d} \quad (1.33)$$

where a , b , c , and d are constants. Solving (1.33) for z , we have

$$z = \frac{-dw + b}{cw - a} \quad (1.34)$$

This shows that when w is a bilinear transformation of z , z is also a bilinear

transformation of w . From (1.27) and (1.28), Z/Z_0 and r are bilinear transformations of each other. In this section, we shall discuss some of elementary properties of bilinear transformations and derive an equivalent circuit of two-port networks. Suppose that v is a bilinear transformation of w , i.e.,

$$v = \frac{ew + f}{gw + h} \quad (1.35)$$

where e, f, g , and h are constants. Since the substitution of (1.33) into (1.35) gives

$$v = \frac{(ea + fc)z + (eb + fd)}{(ga + hc)z + (gb + hd)}$$

v is a bilinear transformation of z . It follows from this that a bilinear transformation of a bilinear transformation is also a bilinear transformation.

Next, let us prove the following important theorem of bilinear transformations: Any circle on the complex plane is transformed into a circle by a bilinear transformation. In this statement, a circle can be a straight line as the limiting case of increasing the radius and removing the center away from the origin. First, let us study a particular case in which the bilinear transformation w of z is simply the inverse of z and the circle on the z plane passes through the origin. Let D be the diameter and ϕ be the angle between the real axis and the diameter through the origin, as shown in Fig. 1.13.

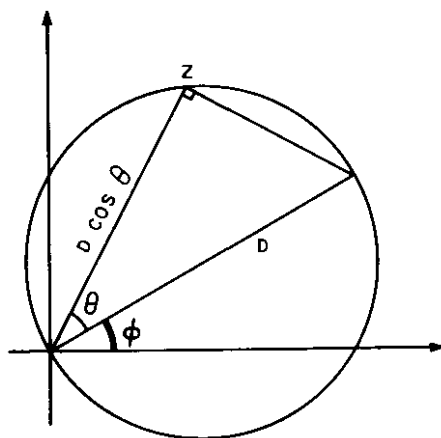


Fig. 1.13. Circle passing through the origin.

Then, the circle is represented by

$$z = D \cos \theta e^{j(\varphi + \theta)}$$

where θ varies from $-\pi/2$ to $\pi/2$. Taking the inverse,

$$\frac{1}{z} = \frac{e^{-j\varphi} e^{-j\theta}}{D \cos \theta} = \frac{e^{-j\varphi} \cos \theta - j \sin \theta}{D \cos \theta} = \frac{e^{-j\varphi}}{D} (1 - j \tan \theta)$$

When θ varies, $(1 - j \tan \theta)/D$ represents a straight line parallel to the imaginary axis and $1/D$ away from it. By rotating this straight line around the origin by $-\varphi$, the locus of $1/z$ is obtained. Thus, we see that the inverse image of a circle passing through the origin is a straight line. Conversely, the inverse image of a straight line which does not pass through the origin is a circle passing through the origin. The circles K_2 in Fig. 1.8(b) and C_2 in Fig. 1.9(b) are examples.

Since the inverse of a complex number is obtained by taking the inverse of its magnitude and changing the sign of its angle, the inverse image of a straight line passing through the origin is a straight line symmetric to it with respect to the real axis.

Next, let us consider the case in which the circle does not pass through the origin; $1/z$ can always be expressed in the form

$$\frac{1}{z} = \frac{(1/k)}{-k/(k+z) + 1} - \frac{1}{k} \quad (1.36)$$

where k is a nonzero constant. This is easily proved by simplifying the right-hand side of the equation. When z represents a circle or a straight line which does not pass through the origin, $k+z$ can be made to pass through the origin by choosing a proper nonzero constant k . Then, from the previous discussions, $1/(k+z)$ becomes a straight line. A constant times a straight line is also a straight line, and the addition of a constant does not change the shape of the locus. Therefore, $-k/(k+z) + 1$ is a straight line. This line does not pass through the origin for the following reason. Assume to the contrary, then, there should be a point z which satisfies

$$-k/(k+z) + 1 = z/(k+z) = 0$$

This, however, requires that $z=0$, i.e., the original locus passes through the origin contradicting the hypothesis. Thus, the straight line we obtained does not pass through the origin. The inverse of this line is a circle. Thus, the first term on the right-hand side of (1.36) is a circle and hence $1/z$ is also a

circle. In other words, the inverse image of a circle or a straight line which does not pass through the origin is a circle. Combining this with the results obtained previously, we conclude that the inverse image of a circle or a straight line is a circle or a straight line depending on whether the original locus passes through the origin or not.

We are now in a position to discuss the general case. In general, a bilinear transformation can be rewritten in the form

$$w = \frac{az + b}{cz + d} = \frac{a}{c} + \frac{b - (ad/c)}{cz + d} \quad (1.37)$$

where c is assumed not to be equal to zero. When $c = 0$, w is a simple linear function of z , and it is obvious that w represents a circle when z is a circle. Therefore, we shall concentrate on the case in which $c \neq 0$. When z represents a circle or a straight line on the complex plane, $cz + d$ also represents a circle or a straight line. Since the second term on the right-hand side of (1.37) is a constant times the inverse of $cz + d$, it is a circle or a straight line according to the previous discussion. Adding a/c , w is seen to be a circle or a straight line. This completes the proof of the theorem.

When a circle is transformed into a circle by a bilinear transformation, there are two possible cases: (i) The inside of one circle is transformed to the inside of the other; and (ii) The inside of one circle corresponds to the outside of the other. In the first case, the rotational directions of the

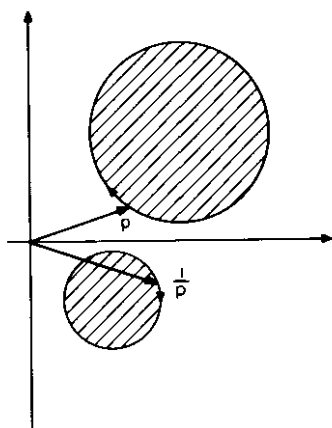


Fig. 1.14. The directions of rotation are the same when the inside of one circle corresponds to the inside of the other.

corresponding points on the circles are the same while in the second case, they are opposite to each other. The proof is as follows: Since any bilinear transformation can be decomposed into the inverse transformation, multiplications by a constant and additions of a constant, and since the last two operations do not change the rotational direction, we have only to discuss the inverse transformations. When the origin is outside of a circle as shown in Fig. 1.14, the transformed circle does not include the origin, and the rotational directions of the corresponding points on the circles are the same. This is obvious from the fact that the inverse of a complex number is represented by a complex vector whose magnitude is the inverse of the original magnitude and whose angle is the negative of the original angle. When the origin is inside of a circle as shown in Fig. 1.15, the trans-

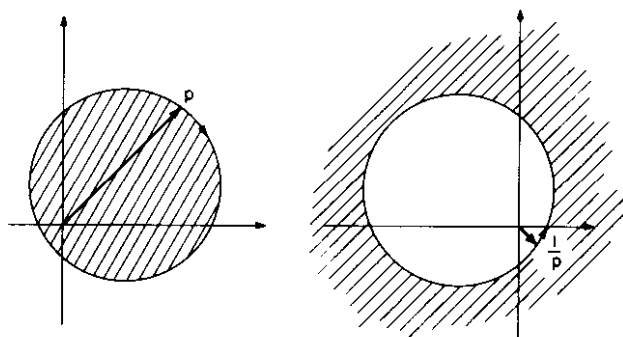


Fig. 1.15. The directions of rotation are opposite when the inside of one circle corresponds to the outside of the other.

formed circle encloses the origin, and, since the inverse of the origin corresponds to infinity, it is obvious that the inside of one circle corresponds to the outside of the other. Furthermore, for the same reason used above, the rotational directions of the corresponding points on the circles become opposite. This completes the proof. From Figs. 1.14 and 1.15, it can be seen that corresponding points move in such a way that corresponding areas would appear on the same side if a person were to stand on each circle facing in the direction of the moving point. This statement is also true when one circle becomes a straight line as a limiting case.

Next, let us derive an equivalent circuit of reciprocal two-port networks as an example of using bilinear transformations. Suppose the relation between the terminal voltages and currents of a given two-port network is

expressed by

$$V_1 = Z_{11}I_1 + Z_{12}I_2 \quad (1.38)$$

$$V_2 = Z_{12}I_1 + Z_{22}I_2 \quad (1.39)$$

where the polarities of the voltages and currents are as shown in Fig. 1.16. Because of the reciprocity of the network under consideration, the coefficient of I_2 in (1.38) is equal to that of I_1 in (1.39). We shall discuss the reciprocity of microwave circuits in detail in Chapter 5.

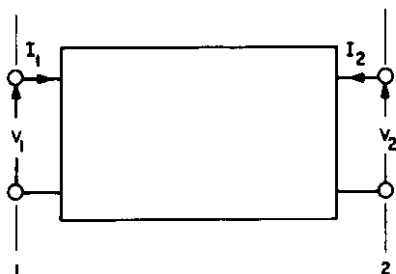


Fig. 1.16. Two-port network.

If a load impedance Z_L is connected to port 2 of the two-port network, V_2 is given by

$$V_2 = -Z_L I_2$$

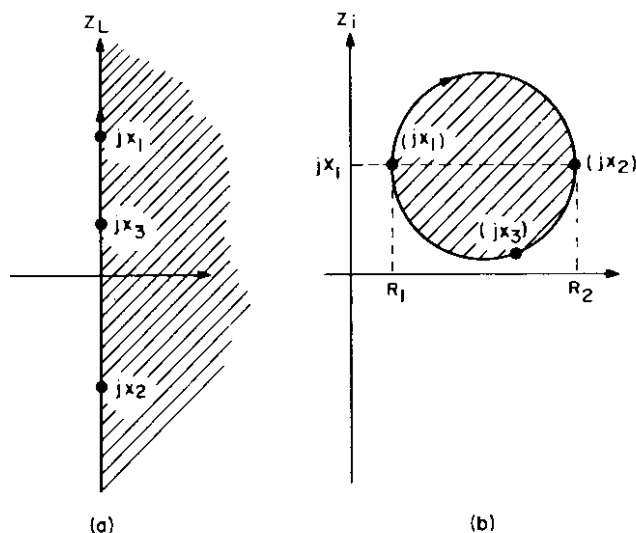
Substituting this into (1.39), I_2 can be obtained in terms of I_1 . Then, substituting this result into (1.38), we have

$$V_1 = \{Z_{11} - Z_{12}^2(Z_L + Z_{22})^{-1}\} I_1$$

This means that when port 2 is terminated by Z_L , the input impedance at port 1 becomes

$$Z_i = V_1/I_1 = \{Z_{11} - Z_{12}^2(Z_L + Z_{22})^{-1}\} \quad (1.40)$$

Thus, Z_i is a bilinear transformation of Z_L . When Z_L changes from $-j\infty$ to $j\infty$ as shown in Fig. 1.17(a), the locus of Z_i is a circle as shown in Fig. 1.17(b). The circle is in the right-half plane with its internal area corresponding to the right-half plane of Z_L for the following reason. Assume the contrary, then Z_i can have a negative real part when Z_L has a positive real part. However, this is impossible since the passive two-port network terminated by a load impedance with a positive real part is a passive circuit and the input impedance must have a positive real part.

Fig. 1.17. Correspondence between Z_L and Z_i .

Let $Z_i(1) = R_1 + jX_1$ be the point with the smallest real part on the circle and let jx_1 be the corresponding Z_L . Similarly, let $Z_i(2) = R_2 + jX_2$ be the point with the largest real part and let jx_2 be the corresponding Z_L . Finally, let jx_3 be a different point on the imaginary axis of Z_L and let the corresponding point on Z_i be as shown by (jx_3) in Fig. 1.17(b). When Z_i traces out a circle, $Z_i - (R_1 + jX_1)$ represents a circle passing through the origin tangential to the imaginary axis as shown in Fig. 1.18(a). The other point at which the circle intersects the real axis is given by $R_2 - R_1$. The inverse of this circle, i.e. the admittance, becomes a constant conductance line with a value $1/(R_2 - R_1)$ as shown in Fig. 1.18(b). When $1/(R_2 - R_1)$ is subtracted, the admittance becomes a pure susceptance, jB . It is easily seen from Figs. 1.17 and 1.18 that the susceptance becomes zero when Z_L is equal to jx_2 and infinite when Z_L is equal to jx_1 . An equivalent circuit for jB with these properties is shown in Fig. 1.19. The transformer ratio n is to be determined by the requirement that jB becomes jB_3 , shown in Fig. 1.18(b), when Z_L becomes jx_3 . Since the admittance is increased n^2 times by the transformer, it follows that

$$[\{j(x_3 - x_1)\}^{-1} + \{j(x_1 - x_2)\}^{-1}] n^2 = jB_3$$

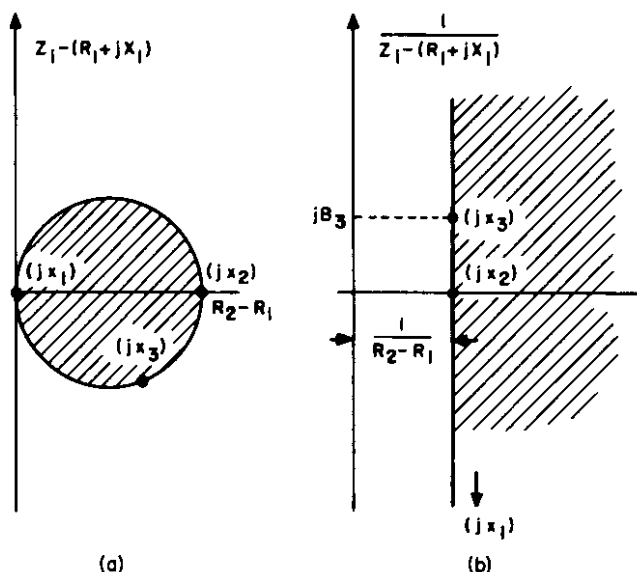
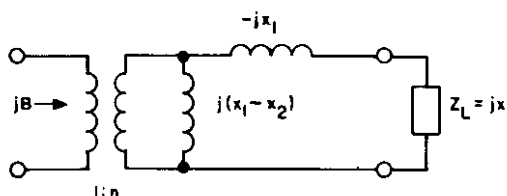


Fig. 1.18. Complex transformations.

Fig. 1.19. An equivalent circuit for $jB = \{Z_i - (R_1 + jX_1)\}^{-1} - \{R_2 - R_1\}^{-1}$.

or equivalently,

$$n^2 = \frac{-B_3(x_3 - x_1)(x_1 - x_2)}{x_3 - x_2} \quad (1.41)$$

When $Z_i(3)$, shown by (jx_3) in Fig. 1.17(b), is located below the jX_1 level, B_3 is positive as shown in Fig. 1.18(b).

From the discussion on the directions of corresponding moving points, it is clear that the point on the circle rotates clockwise when the corresponding point moves upward on the imaginary axis of the Z_L plane. Therefore, only one of three cases is possible, depending on the location of the point corresponding to $Z_L = j\infty$ on the Z_i plane $x_1 > x_3 > x_2$, $x_3 > x_2 > x_1$, $x_2 > x_1 > x_3$. For each of these three cases, since $B_3 > 0$, Eq. (1.41) gives positive n^2 and, hence, the transformer ratio n is realizable. Similarly, if the

point (jx_3) is located above the jX_1 level, B_3 is negative and one of three inequalities,

$$x_2 > x_3 > x_1, \quad x_3 > x_1 > x_2, \quad x_1 > x_2 > x_3$$

has to be satisfied. This means that n^2 is again positive and the transformer ratio is realizable. Thus, the circuit shown in Fig. 1.19 is realizable and gives correct values of jB at three different values of $Z_L = jx$.

It follows from the foregoing discussion that the circuit shown in Fig. 1.20 gives the correct values of input impedance $Z_i(1)$, $Z_i(2)$, and $Z_i(3)$ for three different values of Z_L , jx_1 , jx_2 , and jx_3 . Note that the impedance $R_1 + jX_1$ and conductance $1/(R_2 - R_1)$ subtracted out to obtain jB are all

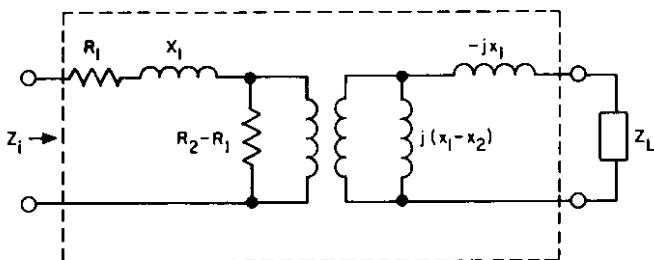


Fig. 1.20. An equivalent circuit for reciprocal two-port networks.

restored in Fig. 1.20. Let us consider the two-port network inside the dotted line. This can be specified by equations similar to (1.38) and (1.39). Let Z'_{11} , Z'_{12} , and Z'_{22} be the coefficients corresponding to Z_{11} , Z_{12} , and Z_{22} , respectively. Then, we have

$$Z_i(k) = \{Z'_{11} - Z'^2_{12}(jx_k + Z'_{22})^{-1}\} \quad (k = 1, 2, 3)$$

Multiplying both sides by $jx_k + Z'_{22}$, a little manipulation gives

$$jx_k Z'_{11} - Z_i(k) Z'_{22} + (Z'_{11} Z'_{22} - Z'^2_{12}) = Z_i(k) jx_k \quad (k = 1, 2, 3) \quad (1.42)$$

This can be considered as three simultaneous linear equations for three unknown quantities Z'_{11} , Z'_{22} , and $Z'_{11} Z'_{22} - Z'^2_{12}$. When the equations are solved, these quantities are uniquely determined in terms of x_k and $Z_i(k)$. On the other hand, from the network defined by (1.38) and (1.39), three simultaneous equations are also obtained for Z_{11} , Z_{22} , and $Z_{11} Z_{22} - Z^2_{12}$ with identical coefficients to those in (1.42). Therefore, Z'_{11} , Z'_{22} , and $Z'_{11} Z'_{22} - Z'^2_{12}$ must be equal to Z_{11} , Z_{22} , and $Z_{11} Z_{22} - Z^2_{12}$, respectively.

That is,

$$Z_{11} = Z'_{11}, \quad Z_{22} = Z'_{22}, \quad Z_{12}^2 = Z_{12}'^2 \quad (1.43)$$

From (1.43) we see, except for the ambiguity in the sign of Z'_{12} , that Fig. 1.20 is equivalent to the given network. The ambiguity in the sign of Z'_{12} can not be eliminated from the above discussion. However, since n^2 but not n itself is specified by (1.41), if the sign is opposite to that of the given network, Z'_{12} can be made equal to Z_{12} by reversing the transformer polarity without disturbing other conditions. Thus, we conclude that the desired equivalent circuit is given by Fig. 1.20 with a proper choice of transformer polarity.

When the network is lossless, the resistances in the equivalent circuit can be eliminated. Furthermore, if a certain length of transmission line is considered part of the network, the equivalent circuit takes a particularly simple form. Suppose that a transmission line is connected at port 1 of the given network and that port 2 is open-circuited. Then, since the circuit is lossless, the incident and reflected waves $a(z)$ and $b(z)$ must satisfy the conditions $|a(z)|^2 - |b(z)|^2 = 0$, $|r| = 1$, and the standing wave ratio is infinite. Let us shift the reference plane (along the transmission line) to a voltage maximum point and consider the circuit consisting of the given two-port network and transmission line beyond this new reference plane, as shown by Fig. 1.21.

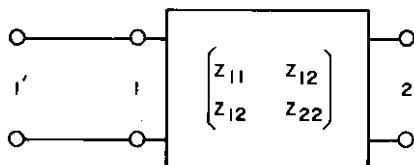


Fig. 1.21. Composite circuit of two-port network and a length of transmission line.

For this new two-port network, $Z_i = \infty$ when $Z_L = \infty$. Therefore, the parallel reactance $j(x_1 - x_2)$ in Fig. 1.20 is also eliminated. The equivalent circuit becomes a series reactance, a transformer, and another series reactance. However, the series reactance on the left of the transformer can be transferred to the right after multiplying by n^2 . Therefore, the equivalent circuit of Fig. 1.21 becomes simply a transformer and a series reactance as shown in Fig. 1.22. Note that in addition to these two parameters, the length of the transmission line has been specified thus making a total of three independent parameters. To determine the transformer ratio n and the reactance X , the Z_i 's for two different Z_L 's other than $Z_L = \infty$ are necessary.

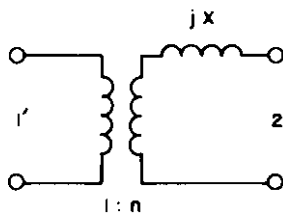


Fig. 1.22. An equivalent circuit of Fig. 1.21.

For example, let Z_i be zero when $Z_L = -jx_1$, then X is equal to x_1 . Furthermore, if $Z_i = jX_0$ when $Z_L = 0$, n^2 is given by X/X_0 .

1.4 Power Waves

Let us consider a linear one-port network containing voltage and current sources. First suppose that a voltage source with the same magnitude as that of the open-circuited voltage at the terminals, but with the opposite phase, is inserted in series before connecting to a load. The effects of all the sources within the network and the added voltage source cancel each other and the circuit as a whole acts as a simple impedance Z_g as far as the outside phenomena are concerned. Next, suppose that the sources inside the one-port network are all nullified and the phase of the added voltage source is reversed. The super-position of these two cases is equivalent to the original one-port network. As a result, any linear one-port network can be represented by the series connection of Z_g and a voltage source equal to the open-circuited terminal voltage. Therefore, in the following discussion we shall use Fig. 1.23 to represent an arbitrary linear one-port generator. Referring

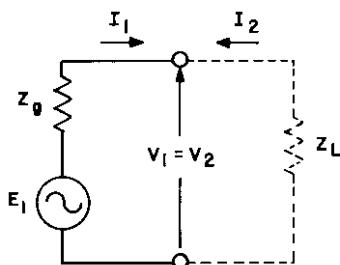


Fig. 1.23. An equivalent circuit of one-port generators.

to Fig. 1.23, let V_L be the voltage across the load impedance Z_L and I_L be the current flowing through it. The power consumed in the load is given by

$$P_L = \operatorname{Re} \{Z_L\} |I_L|^2$$

Since the magnitude of the current is $|E_1/(Z_L + Z_g)|$, P_L is rewritten in the form

$$P_L = R_L \left| \frac{E_1}{Z_L + Z_g} \right|^2 = \frac{R_L |E_1|^2}{(R_L + R_g)^2 + (X_L + X_g)^2} \quad (1.44)$$

$$= |E_1|^2 / 4R_g + \frac{(R_L - R_g)^2}{R_L} + \frac{(X_L + X_g)^2}{R_L} \quad (1.45)$$

where R_L and R_g are the real parts of Z_L and Z_g , respectively and X_L and X_g are the imaginary parts. With $R_g > 0$, we can easily see from (1.45), that the power becomes maximum when

$$R_L = R_g, \quad X_L = -X_g \quad (1.46)$$

This maximum power is called the available power of the generator and indicated by P_a :

$$P_a = |E_1|^2 / 4R_g \quad (R_g > 0) \quad (1.47)$$

and P_a is solely determined by the parameters of the generator, hence it is a characteristic of the generator independent of the load. In order to realize the maximum power consumption, (1.46) has to be satisfied. This condition, $Z_L = Z_g^*$, is called the matching condition and if it is not satisfied, the load consumes less power than P_a .

When R_g is negative, P_L becomes infinite as R_L and X_L approach $-R_g$ and $-X_g$, respectively, as can easily be seen from (1.44), therefore the available power is infinite if R_g is negative. However, the power expressed by the right-hand side of (1.47) remains finite for any nonzero R_g , and it is called the exchangeable power P_e of the generator. That is,

$$P_e = |E_1|^2 / 4R_g \quad (R_g \leq 0) \quad (1.48)$$

When R_g is positive, P_e is equal to the available power from the generator, however, if R_g is negative, it is no longer the maximum power available. Since a small variation in Z_L from Z_g^* produces only a second order variation in the right-hand side of (1.45), P_e can be considered as the stationary value of P_L with respect to a variation of Z_L . Sometimes, P_e is called the characteristic power of the generator, emphasizing that it is

invariant to nonsingular lossless transformations as we shall explain in Section 5.3 and again in Section 7.1.

With this preparation, let us introduce power waves a_1 and b_1 which are defined by linear transformations of V_1 and I_1 :

$$a_1 = \frac{1}{2} |\operatorname{Re} Z_g|^{-1/2} (V_1 + Z_g I_1), \quad b_1 = \frac{1}{2} |\operatorname{Re} Z_g|^{-1/2} (V_1 - Z_g^* I_1) \quad (1.49)$$

Note the similarity between these equations and the waves (1.16) introduced in Section 1.1. With a fixed Z_g , and if V_1 and I_1 are given, then a_1 and b_1 can be readily calculated from (1.49). On the other hand, if a_1 and b_1 are given, V_1 and I_1 are obtained from the inverse transformation

$$V_1 = p_1 |\operatorname{Re} Z_g|^{-1/2} (Z_g^* a_1 + Z_g b_1), \quad I_1 = p_1 |\operatorname{Re} Z_g|^{-1/2} (a_1 - b_1) \quad (1.50)$$

where p_1 is defined by

$$p_1 = \begin{cases} 1 & \text{when } \operatorname{Re} Z_g > 0 \\ -1 & \text{when } \operatorname{Re} Z_g < 0 \end{cases} \quad (1.51)$$

Thus, any result in terms of one set of variables can easily be converted to that in terms of the other set of variables. This justifies the use of a_1 and b_1 defined by (1.49) in place of the terminal voltage and current for any analysis. Referring to Fig. 1.23, V_1 is given by

$$V_1 = E_1 - Z_g I_1$$

From this and (1.49), we have

$$|a_1|^2 = |E_1|^2 / 4 |R_g|$$

which is equivalent to

$$P_e = |a_1|^2 |R_g| / R_g = p_1 |a_1|^2 \quad (1.52)$$

Next, let us consider the meaning of $|a_1|^2 - |b_1|^2$. From (1.49), this becomes

$$\begin{aligned} |a_1|^2 - |b_1|^2 &= \frac{(V_1 + Z_g I_1)(V_1^* + Z_g^* I_1^*) - (V_1 - Z_g^* I_1)(V_1^* - Z_g I_1^*)}{4 |R_g|} \\ &= \frac{(Z_g + Z_g^*)(V_1 I_1^* + V_1^* I_1)}{4 |R_g|} \\ &= \frac{R_g}{|R_g|} \operatorname{Re} \{V_1 I_1^*\} \end{aligned}$$

from which we have

$$\operatorname{Re} \{V_1 I_1^*\} = p_1 (|a_1|^2 - |b_1|^2) \quad (1.53)$$

Since $\operatorname{Re} \{V_1 I_1^*\}$ is the net power transferred from the generator to the

load, we see from (1.52) and (1.53) that the exchangeable power is given by $p_1|a_1|^2$ and the net power by $p_1(|a_1|^2 - |b_1|^2)$. This leads to the following interpretation: The generator sends the exchangeable power $p_1|a_1|^2$ to the load, however, $p_1|b_1|^2$ is reflected back to the generator, and hence the net power to the load is given by $p_1(|a_1|^2 - |b_1|^2)$, where a_1 and b_1 are the waves associated with these forward and reflected powers. Since $|b_1|^2$ is nonnegative, $|a_1|^2 - |b_1|^2$ becomes maximum when $|b_1|^2$ is equal to zero. Therefore, whether the load contains some power sources or not, the magnitude of the exchangeable power can be identified as the maximum power that the generator can supply when $R_g > 0$ and the maximum power that the generator can absorb when $R_g < 0$. In other words, the maximum power that the generator can exchange with the external circuit is the exchangeable power although the generator can absorb, or supply, more than $|P_e|$ when $R_g > 0$ or $R_g < 0$, respectively.

Power waves a_1 and b_1 use the generator impedance in their definition. Similar waves a_2 and b_2 can be defined from the other side of the center line of the circuit in Fig. 1.23, using the load impedance rather than generator impedance in the definition. In this case, the direction of the current has to be reversed. Thus, we have

$$\begin{aligned} a_2 &= \frac{1}{2} |\operatorname{Re} Z_L|^{-1/2} (V_2 + Z_L I_2) = \frac{1}{2} |\operatorname{Re} Z_L|^{-1/2} (V_1 - Z_L I_1) \\ b_2 &= \frac{1}{2} |\operatorname{Re} Z_L|^{-1/2} (V_2 - Z_L^* I_2) = \frac{1}{2} |\operatorname{Re} Z_L|^{-1/2} (V_1 + Z_L^* I_1) \end{aligned} \quad (1.54)$$

Note that a_1 and b_1 are not necessarily equal to b_2 and a_2 , respectively. Let p_2 be +1 or -1 when the sign of $\operatorname{Re} Z_L$ is positive or negative, respectively. Then, $p_2|a_2|^2$ represents the exchangeable power of the load. If the open-circuited terminal voltage of the load is zero, this is equal to zero. The net power transferred from the load to the generator is given by $p_2(|a_2|^2 - |b_2|^2)$. Therefore, the net power from the generator to the load is given by $p_2(|b_2|^2 - |a_2|^2)$. Although $p_1|b_1|^2$ was considered as the power reflected from the load back to the generator, we shall call $p_2|b_2|^2$ the actual power flowing into the load. Note that because of the voltage source in the generator, it can be a finite quantity even when the exchangeable power of the load is zero and hence the load does not send out any power to be reflected back. The actual power is equal to the net power transferred to the load plus the exchangeable power of the load, if any. Of course, $p_1|b_1|^2$ could be called the actual power flowing into the generator.

Let us next introduce the power wave reflection coefficient s_1 defined by

$$s_1 = (b_1/a_1) \quad (1.55)$$

When the exchangeable power of the load is equal to zero, $V_1 = Z_L I_1$ and s_1 can be expressed in terms of impedances:

$$s_1 = (Z_L - Z_g^*) / (Z_L + Z_g) \quad (1.56)$$

Substituting $Z_g = R_g + jX_g$ and $Z_L = R_L + jX_L$ into (1.56), s_1 can be rewritten in the form

$$s_1 = \frac{R_L + j(X_L + X_g) - R_g}{R_L + j(X_L + X_g) + R_g} \quad (1.57)$$

Comparing this expression with that of the reflection coefficient (1.28), we see that s_1 corresponds to the vector drawn from the center of the Smith chart to the point where the normalized impedance is given by $\{R_L + j(X_L + X_g)\} / R_g$. In other words, if the reactive part of Z_g is added to Z_L and normalized with respect to the real part of Z_g , the corresponding point on the Smith chart gives the magnitude and phase of the power wave reflection coefficient. From this, the following important property of s_1 is derived: When R_g and R_L have the same sign, $|s_1| < 1$, and when they have opposite signs, $|s_1| > 1$.

The power reflection coefficient is given by

$$|s_1|^2 = |(Z_L - Z_g^*) / (Z_L + Z_g)|^2 \quad (1.58)$$

When the matching condition (1.46) is satisfied, the power reflection coefficient becomes zero, as is expected.

Looking into the load from the generator, s_1 and $|s_1|^2$ are the reflection coefficients. The corresponding reflection coefficients s_2 and $|s_2|^2$ looking into the generator from the load are given by

$$s_2 = (Z_g - Z_L^*) / (Z_g + Z_L), \quad |s_2|^2 = |(Z_g - Z_L^*) / (Z_g + Z_L)|^2$$

where the subscripts g and L are interchanged. The reflection coefficient s_2 is not necessarily equal to s_1 . However, since $|Z_g - Z_L^*| = |Z_g^* - Z_L| = |Z_L - Z_g^*|$, $|s_2|^2$ is always equal to $|s_1|^2$.

The quantity $1 - |s_1|^2$ is called the power transmission coefficient which is, of course, equal to $1 - |s_2|^2$. It is worth noting that the power transmission coefficient multiplied by the exchangeable power of the generator is equal to the net power transferred to the load and conversely the net power divided by the power transmission coefficient is the exchangeable power.

If Z_g is replaced by positive real Z_0 , the power waves a_1 and b_1 defined

by (1.49) become identical to the traveling waves $a(z)$ and $b(z)$ given by (1.16). As a result, the various relations derived for power waves also hold for traveling waves provided that Z_g is replaced by positive real Z_0 . One example is the result that the power transmitted in the z -positive direction is given by $|a(z)|^2 - |b(z)|^2$. The same relation is easily derived from (1.53) since p_1 is equal to plus one in that case. On the other hand, when the characteristic impedance is complex, the situation becomes slightly different, and $|a(z)|^2 - |b(z)|^2$ no longer gives the net power $\text{Re}\{VI^*\}$. Furthermore, as we shall see in Chapter 2, when there is no reflection in terms of power waves, the traveling wave has some reflection from the load and vice versa.

PROBLEMS

- 1.1 Derive the transmission line equations corresponding to (1.3) and (1.4) and (1.5) taking into account the resistance of the wire. Then calculate the characteristic impedance and propagation constant of the line.
- 1.2 Suppose that a length of $50\ \Omega$ transmission line is terminated by a $20\ \Omega$ resistor and $30\ \text{pF}$ capacitor connected in series and calculate the SWR and first voltage minimum point at $1\ \text{GHz}$ using the Smith chart.
- 1.3 Suppose that the SWR is 2.6 and the first voltage minimum point is 0.338 wavelengths away from a load terminating a $50\ \Omega$ transmission line. Calculate the impedance of the load on the Smith chart.
- 1.4 Suppose that a $50\ \Omega$ transmission line is terminated by $(15 + j15)\ \Omega$ and the incident power is $5\ \text{kW}$. Calculate the net power consumed in the termination. Also calculate the maximum voltage on the line.
- 1.5 Calculate the inverse of $1.0 + j1.0$ using the Smith chart.
- 1.6 Convert the circuit shown in Fig. 1.24 to the equivalent circuit corresponding to Fig. 1.20.

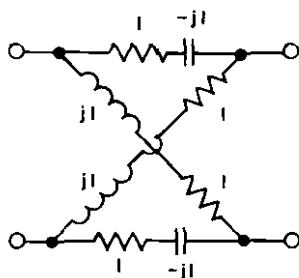


Fig. 1.24. A two-port network.

- 1.7 Prove that the exchangeable power is invariant to the insertion of ideal transformers or the connection of shunt and series reactances.

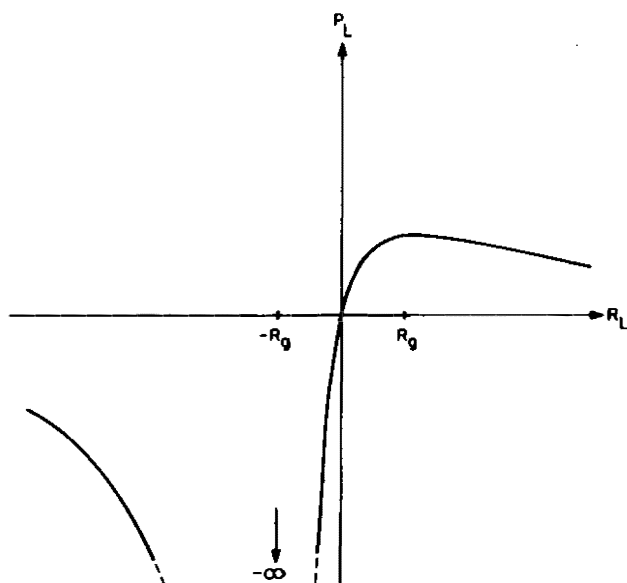


Fig. 1.25. P_L as a function of R_L when $R_g > 0$ and $X_g + X_L = 0$.

- 1.8 When $X_L + X_g$ is equal to zero and $R_g > 0$, P_L varies with R_L as shown in Fig. 1.25. Investigate the case in which $R_g < 0$ and $X_L + X_g$ is not necessarily equal to zero.



CHAPTER 2

ELECTROMAGNETIC FIELD VECTORS

In order to discuss microwave circuits, Maxwell's equations must be studied since these describe the relations between electric and magnetic fields. If the concept of a vector is introduced, the relations between the fields become simpler to describe and easier to understand. Therefore, we shall make extensive use of vectors in this book. This chapter reviews some of the important theorems on vector analysis which will facilitate our later study. In Section 2.1, elements of vectors and vector analysis are presented, and in Section 2.2, Maxwell's equations are explained in terms of the field vectors in order to refresh the reader's understanding of electromagnetic theory. Section 2.3 gives an analysis of plane waves, first in terms of scalar quantities resolving the field vectors into their components, and then in terms of vectors after which the results are compared.

2.1 Vectors

A vector is a quantity with magnitude and direction. To visualize it, we usually consider an arrow whose length and direction expresses the magnitude and direction of the vector, respectively, as shown in Fig. 2.1. The space of



Fig. 2.1. Arrow representing vector A.

this figure is not the actual space in which we live, but rather it is an abstract space where the magnitude and direction of the vector is represented by the length and direction of the corresponding arrow. Vectors are generally functions of time and position existing in actual space, therefore, depending on the position and time, the lengths and the directions of the vector arrows vary in abstract space. Sometimes however, the tail of the arrow is located at the point in the actual space where the vector is considered. An example will be shown in Fig. 2.10. In such a case, the illustrated space has two meanings; one is the actual space, and the other is the abstract space in which we consider the arrows representing vectors. Vectors representing the velocity of liquid, electric field, or magnetic field, which are functions of position, are sometimes represented by a cluster of curves with arrow heads. In this case, the direction of the vector at each point on a curve is given by the tangent in the direction indicated by the arrow. The magnitude of the vector at a particular point is given by the density of curves in the vicinity of that point.

When the length of an arrow shrinks to zero, its direction loses meaning. The corresponding vector is called a zero vector and is indicated by 0 . The vector which has the same direction as a vector $\mathbf{A}$, but with a magnitude k times as large, is indicated by $k\mathbf{A}$. The addition of two vectors $\mathbf{A}$ and $\mathbf{B}$ is defined as follows. Place the tails of the corresponding arrows at the same point in abstract space as shown in Fig. 2.2 and consider the plane which

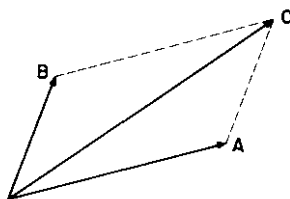


Fig. 2.2. Addition of two vectors $\mathbf{A}$ and $\mathbf{B}$.

they define. On this plane construct a parallelogram with two of its sides coinciding with the arrows. Then, the arrow $\mathbf{C}$ corresponding to the diagonal as shown in Fig. 2.2 represents the vector $\mathbf{A} + \mathbf{B}$. Alternatively, translate $\mathbf{B}$ until its tail coincides with the tip of vector $\mathbf{A}$. If an arrow $\mathbf{C}$ is now drawn from the tail of $\mathbf{A}$ to the tip of $\mathbf{B}$, then $\mathbf{C}$ represents $\mathbf{A} + \mathbf{B}$. In these operations, it is immaterial whether $\mathbf{A}$ or $\mathbf{B}$ comes first, i.e., the addition is commutative. Thus, we have

$$\mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A} = \mathbf{C} \quad (2.1)$$

Furthermore, note that $\mathbf{A}$ and $\mathbf{B}$ could be vectors defined at different times and different points in the actual space.

Vectors with magnitudes equal to unity are called unit vectors. Let the unit vectors $\mathbf{i}_x$, $\mathbf{i}_y$, and $\mathbf{i}_z$ have the directions of the x , y , and z axes of a rectangular coordinate system. Then, as is easily seen from Fig. 2.3, an arbitrary vector $\mathbf{A}$ can be expressed as

$$\mathbf{A} = \mathbf{i}_x A_x + \mathbf{i}_y A_y + \mathbf{i}_z A_z \quad (2.2)$$

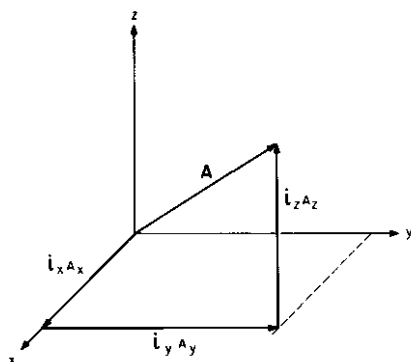


Fig. 2.3. Vector $\mathbf{A}$ and its rectangular components.

A_x , A_y , and A_z are called the x , y , and z components of vector $\mathbf{A}$ while $\mathbf{i}_x A_x$ is called the projection of $\mathbf{A}$ in the x -direction or the projection of $\mathbf{A}$ on $\mathbf{i}_x$.

The scalar product $\mathbf{A} \cdot \mathbf{B}$ of two vectors $\mathbf{A}$ and $\mathbf{B}$ is defined as $|\mathbf{A}| |\mathbf{B}| \cos \theta$ where $|\mathbf{A}|$ and $|\mathbf{B}|$ represent the magnitudes of vectors $\mathbf{A}$ and $\mathbf{B}$, respectively, and θ is the angle between them. Whether $\mathbf{A}$ or $\mathbf{B}$ comes first, the scalar product remains the same:

$$\mathbf{A} \cdot \mathbf{B} = \mathbf{B} \cdot \mathbf{A} = |\mathbf{A}| |\mathbf{B}| \cos \theta \quad (2.3)$$

The scalar product can be considered as the product of the length of $\mathbf{A}$ and the length of $\mathbf{B}'$, where $\mathbf{B}'$ is the projection of $\mathbf{B}$ on $\mathbf{A}$ as shown in Fig. 2.4.

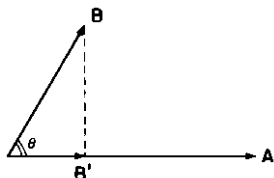


Fig. 2.4. Explanation of scalar product.

If θ is greater than 90° , the value of the scalar product becomes negative.

The scalar product of $\mathbf{A}$ and the addition of two vectors $\mathbf{B}$ and $\mathbf{C}$ is equal to the addition of the scalar product of $\mathbf{A}$ and $\mathbf{B}$ plus the scalar product of $\mathbf{A}$ and $\mathbf{C}$. That is,

$$\mathbf{A} \cdot (\mathbf{B} + \mathbf{C}) = \mathbf{A} \cdot \mathbf{B} + \mathbf{A} \cdot \mathbf{C} \quad (2.4)$$

This is proved as follows. The left-hand side of (2.4) is equal to the length of $\mathbf{A}$ multiplied by the length of the projection $(\mathbf{B} + \mathbf{C})'$ of $(\mathbf{B} + \mathbf{C})$ on $\mathbf{A}$. The right-hand side is equal to the length of $\mathbf{A}$ times the length of the projection $\mathbf{B}'$ of $\mathbf{B}$ on $\mathbf{A}$ plus the length of $\mathbf{A}$ times the length of the projection $\mathbf{C}'$ of $\mathbf{C}$ on $\mathbf{A}$. Referring to Fig. 2.5, however, we see that the projection $(\mathbf{B} + \mathbf{C})'$ is

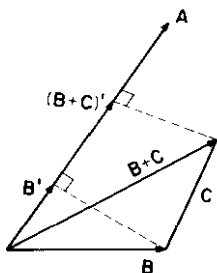


Fig. 2.5. Scalar product of a vector and the sum of two vectors.

equal to the addition of the projection $\mathbf{B}'$ and $\mathbf{C}'$. Thus, we have

$$\begin{aligned} \mathbf{A} \cdot (\mathbf{B} + \mathbf{C}) &= \mathbf{A} \cdot (\mathbf{B} + \mathbf{C})' = |\mathbf{A}| |(\mathbf{B} + \mathbf{C})'| = |\mathbf{A}| \{|\mathbf{B}'| + |\mathbf{C}'|\} \\ &= \mathbf{A} \cdot \mathbf{B}' + \mathbf{A} \cdot \mathbf{C}' = \mathbf{A} \cdot \mathbf{B} + \mathbf{A} \cdot \mathbf{C} \end{aligned} \quad (2.5)$$

This completes the proof. From (2.5) and the relations $\mathbf{i}_x \cdot \mathbf{i}_x = 1$, $\mathbf{i}_x \cdot \mathbf{i}_y = 0$, etc., the scalar product of $\mathbf{A}$ and $\mathbf{B}$ can be expressed in terms of the x , y , z components as follows:

$$\mathbf{A} \cdot \mathbf{B} = (\mathbf{i}_x A_x + \mathbf{i}_y A_y + \mathbf{i}_z A_z) \cdot (\mathbf{i}_x B_x + \mathbf{i}_y B_y + \mathbf{i}_z B_z) = A_x B_x + A_y B_y + A_z B_z \quad (2.6)$$

The vector product $\mathbf{A} \times \mathbf{B}$ of two vectors $\mathbf{A}$ and $\mathbf{B}$ is a vector defined as follows. The magnitude of the vector is equal to the area of parallelogram defined by $\mathbf{A}$ and $\mathbf{B}$ while its direction is normal to the plane of the parallelogram and given by the right-hand screw rule as shown in Fig. 2.6. A screw would advance in the direction of the vector when turned from $\mathbf{A}$ to $\mathbf{B}$ through the smaller angle θ as shown. Referring to Fig. 2.6, the magnitude

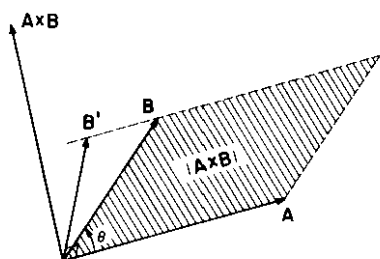


Fig. 2.6. Explanation of vector product.

of $\mathbf{A} \times \mathbf{B}$ is equal to $|\mathbf{A}| |\mathbf{B}| \sin \theta$. The vector product $\mathbf{B} \times \mathbf{A}$ has the same magnitude as $\mathbf{A} \times \mathbf{B}$. However, since the right-handed screw advances in the opposite direction when it is turned from $\mathbf{B}$ to $\mathbf{A}$, the direction of $\mathbf{B} \times \mathbf{A}$ is opposite to the direction of $\mathbf{A} \times \mathbf{B}$. That is,

$$\mathbf{B} \times \mathbf{A} = -\mathbf{A} \times \mathbf{B} \quad (2.7)$$

If we draw a straight line parallel to $\mathbf{A}$ passing through the tip of $\mathbf{B}$ and take an arbitrary vector drawn from the tail of $\mathbf{A}$ to this line, then the vector product of this new vector and $\mathbf{A}$ is equal to $\mathbf{A} \times \mathbf{B}$. Let $\mathbf{B}'$ be one normal to $\mathbf{A}$ as shown in Fig. 2.6. Then, of course, we have

$$\mathbf{A} \times \mathbf{B} = \mathbf{A} \times \mathbf{B}' \quad (2.8)$$

Suppose that one of the vectors in a vector product, say the second one, is expressed as an addition of two vectors $\mathbf{B}$ and $\mathbf{C}$. Then,

$$\mathbf{A} \times (\mathbf{B} + \mathbf{C}) = \mathbf{A} \times \mathbf{B} + \mathbf{A} \times \mathbf{C} \quad (2.9)$$

The proof is as follows. In Fig. 2.7, dotted lines are drawn from the tips of vectors $\mathbf{B}$, $\mathbf{C}$, and $\mathbf{B} + \mathbf{C}$ parallel with $\mathbf{A}$, and the projections of these vectors on to the plane perpendicular to $\mathbf{A}$ are indicated by $\mathbf{B}'$, $\mathbf{C}'$, and $(\mathbf{B} + \mathbf{C})'$, respectively. From (2.8), we have

$$\begin{aligned} \mathbf{A} \times (\mathbf{B} + \mathbf{C}) &= \mathbf{A} \times (\mathbf{B} + \mathbf{C})' \\ \mathbf{A} \times \mathbf{B} &= \mathbf{A} \times \mathbf{B}' \\ \mathbf{A} \times \mathbf{C} &= \mathbf{A} \times \mathbf{C}' \end{aligned} \quad (2.10)$$

These vector products lie in the plane perpendicular to $\mathbf{A}$ since they are all normal to $\mathbf{A}$. They are obtained by rotating $(\mathbf{B} + \mathbf{C})'$, $\mathbf{B}'$, and $\mathbf{C}'$ by 90° around $\mathbf{A}$ and multiplying the lengths by a factor $|\mathbf{A}|$, respectively. From Fig. 2.7, we see that the projection of $\mathbf{B} + \mathbf{C}$ to a plane is equal to the

addition of the projections of **B** and **C** to the same plane, i.e.,

$$(\mathbf{B} + \mathbf{C})' = \mathbf{B}' + \mathbf{C}'$$

Since neither the rotation of each vector in the same direction by the same angle nor the multiplication of its length by the same factor does change the additive relation between the vectors, we have

$$\mathbf{A} \times (\mathbf{B} + \mathbf{C})' = \mathbf{A} \times \mathbf{B}' + \mathbf{A} \times \mathbf{C}'$$

The substitution of (2.10) gives (2.9).

The repeated use of (2.9) and the relations $\mathbf{i}_x \times \mathbf{i}_x = 0$, $\mathbf{i}_x \times \mathbf{i}_y = \mathbf{i}_z$, etc., gives $\mathbf{A} \times \mathbf{B}$ in terms of the rectangular components of **A** and **B** as follows:

$$\begin{aligned} \mathbf{A} \times \mathbf{B} &= (\mathbf{i}_x A_x + \mathbf{i}_y A_y + \mathbf{i}_z A_z) \times (\mathbf{i}_x B_x + \mathbf{i}_y B_y + \mathbf{i}_z B_z) \\ &= \mathbf{i}_x (A_y B_z - A_z B_y) + \mathbf{i}_y (A_z B_x - A_x B_z) + \mathbf{i}_z (A_x B_y - A_y B_x) \end{aligned} \quad (2.11)$$

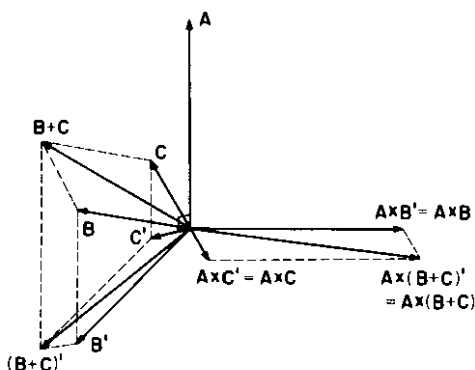


Fig. 2.7. Vector product of a vector and the sum of two vectors.

Treating $\mathbf{i}_x$, $\mathbf{i}_y$, $\mathbf{i}_z$ as if they were ordinary numbers, the above formula can be expressed as the determinant

$$\mathbf{A} \times \mathbf{B} = \begin{vmatrix} \mathbf{i}_x & \mathbf{i}_y & \mathbf{i}_z \\ A_x & A_y & A_z \\ B_x & B_y & B_z \end{vmatrix} \quad (2.12)$$

The vector product of two vectors **B** and **C**, namely, $\mathbf{B} \times \mathbf{C}$, is itself a vector. We can, therefore, consider the scalar product with another vector **A**, i.e., $\mathbf{A} \cdot (\mathbf{B} \times \mathbf{C})$. Referring to Fig. 2.8, $\mathbf{B} \times \mathbf{C}$ has a magnitude which is equal to the area of the parallelogram defined by **B** and **C** and a direction

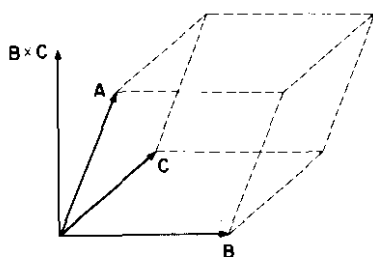


Fig. 2.8. Scalar product of a vector and the vector product of two vectors.

normal to the plane of the parallelogram. Since the magnitude of $\mathbf{A} \cdot (\mathbf{B} \times \mathbf{C})$ is equal to the magnitude of $(\mathbf{B} \times \mathbf{C})$ times the length of the projection of $\mathbf{A}$ on $(\mathbf{B} \times \mathbf{C})$, it is equal to the volume, or the base area times the height, of the parallelepiped defined by $\mathbf{A}$, $\mathbf{B}$, and $\mathbf{C}$. This volume remains the same whether the parallelogram defined by $\mathbf{A}$ and $\mathbf{B}$ or the one defined by $\mathbf{C}$ and $\mathbf{A}$ is considered as the base. Thus, we have

$$\mathbf{A} \cdot (\mathbf{B} \times \mathbf{C}) = \mathbf{C} \cdot (\mathbf{A} \times \mathbf{B}) = \mathbf{B} \cdot (\mathbf{C} \times \mathbf{A}) \quad (2.13)$$

However, if $\mathbf{C} \times \mathbf{B}$ is taken instead of $\mathbf{B} \times \mathbf{C}$, the direction of the vector is opposite and the volume becomes negative.

In a similar way to the preceding discussion, we can consider the vector product of $\mathbf{A}$ and $(\mathbf{B} \times \mathbf{C})$. This is represented by $\mathbf{A} \times (\mathbf{B} \times \mathbf{C})$ where the parentheses indicate that the inside operation must be carried out first. Since it is perpendicular to $\mathbf{B} \times \mathbf{C}$, which is in turn normal to the plane defined by $\mathbf{B}$ and $\mathbf{C}$, $\mathbf{A} \times (\mathbf{B} \times \mathbf{C})$ lies in this plane. This means that $\mathbf{A} \times (\mathbf{B} \times \mathbf{C})$ is expressible as a linear combination of $\mathbf{B}$ and $\mathbf{C}$, i.e., the sum of some coefficients times $\mathbf{B}$ and $\mathbf{C}$. These coefficients are found to be $(\mathbf{A} \cdot \mathbf{C})$ and $-(\mathbf{A} \cdot \mathbf{B})$, respectively. That is,

$$\mathbf{A} \times (\mathbf{B} \times \mathbf{C}) = (\mathbf{A} \cdot \mathbf{C}) \mathbf{B} - (\mathbf{A} \cdot \mathbf{B}) \mathbf{C} \quad (2.14)$$

Since this formula will be used several times, we shall prove it as follows. First consider the special case in which $\mathbf{A}$ is equal to $\mathbf{B}$. In Fig. 2.9, a dotted line parallel to $\mathbf{B}$ is drawn passing through the tip of $\mathbf{C}$, and $\mathbf{C}'$ is drawn perpendicular to $\mathbf{B}$ from the tail of $\mathbf{B}$ to the dotted line. From (2.8), $\mathbf{B} \times \mathbf{C}$ is equal to $\mathbf{B} \times \mathbf{C}'$. The vector $\mathbf{C}' - \mathbf{C}$, drawn from the tip of $\mathbf{C}$ to that of $\mathbf{C}'$, is parallel to $\mathbf{B}$ but the direction is opposite. Its magnitude is equal to that of the projection of $\mathbf{C}$ on $\mathbf{B}$, which is given by $|\mathbf{C}| \cos \theta$ where θ is the angle between $\mathbf{B}$ and $\mathbf{C}$. Since this magnitude can be expressed as $(|\mathbf{B}| |\mathbf{C}| \cos \theta /$

$|\mathbf{B}|^2 |\mathbf{B}|$, $\mathbf{C}' - \mathbf{C}$ is equal to $-\{(\mathbf{B} \cdot \mathbf{C})/(\mathbf{B} \cdot \mathbf{B})\} \mathbf{B}$. From this, we have

$$\mathbf{C}' = \mathbf{C} - \mathbf{B}(\mathbf{B} \cdot \mathbf{C})/(\mathbf{B} \cdot \mathbf{B}) \quad (2.15)$$

Noting that $\mathbf{B} \times (\mathbf{B} \times \mathbf{C}')$ is perpendicular to both $\mathbf{B}$ and $\mathbf{B} \times \mathbf{C}'$, it follows from Fig. 2.9 that the direction of the vector $\mathbf{B} \times (\mathbf{B} \times \mathbf{C}')$ is parallel but opposite to $\mathbf{C}'$ and its magnitude is given by $|\mathbf{B}| (|\mathbf{B}| |\mathbf{C}'|)$. Therefore, we have

$$\mathbf{B} \times (\mathbf{B} \times \mathbf{C}') = -(\mathbf{B} \cdot \mathbf{B}) \mathbf{C}'$$

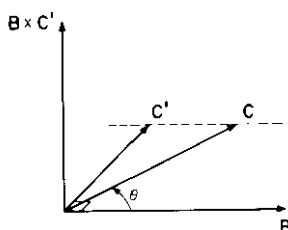


Fig. 2.9. Explanation of vector product formula.

Substituting (2.15) in the right-hand side and using the fact that $\mathbf{B} \times \mathbf{C}'$ is equal to $\mathbf{B} \times \mathbf{C}$, the desired result

$$\mathbf{B} \times (\mathbf{B} \times \mathbf{C}) = (\mathbf{B} \cdot \mathbf{C}) \mathbf{B} - (\mathbf{B} \cdot \mathbf{B}) \mathbf{C} \quad (2.16)$$

is obtained. This shows that (2.14) is correct in this special case where $\mathbf{A} = \mathbf{B}$.

Let us now consider the general case in which $\mathbf{A}$ is not necessarily equal to $\mathbf{B}$. From the explanation preceding (2.14), $\mathbf{A} \times (\mathbf{B} \times \mathbf{C})$ can be written in the form

$$\mathbf{A} \times (\mathbf{B} \times \mathbf{C}) = k_1 \mathbf{B} + k_2 \mathbf{C} \quad (2.17)$$

We must, therefore, determine the coefficients k_1 and k_2 . To do so, first multiply (2.17) by $\mathbf{A} \cdot$ from the left, then we have

$$\mathbf{A} \cdot \mathbf{A} \times (\mathbf{B} \times \mathbf{C}) = k_1 (\mathbf{A} \cdot \mathbf{B}) + k_2 (\mathbf{A} \cdot \mathbf{C})$$

The left-hand side of this equation represents the volume of parallelepiped defined by $\mathbf{A}$, $\mathbf{A}$, and $\mathbf{B} \times \mathbf{C}$. However, since the two edges coincide with each other, the volume is zero. Setting the right-hand side equal to zero, we obtain

$$k_2 = -k_1 (\mathbf{A} \cdot \mathbf{B})/(\mathbf{A} \cdot \mathbf{C})$$

Substituting into (2.17), we have

$$\mathbf{A} \times (\mathbf{B} \times \mathbf{C}) = k_1 \{ \mathbf{B} - \mathbf{C}(\mathbf{A} \cdot \mathbf{B})/(\mathbf{A} \cdot \mathbf{C}) \} \quad (2.18)$$

Multiplying (2.18) by $\mathbf{B} \cdot$, we have

$$\mathbf{B} \cdot \mathbf{A} \times (\mathbf{B} \times \mathbf{C}) = k_1 \{ (\mathbf{B} \cdot \mathbf{B}) - (\mathbf{B} \cdot \mathbf{C})(\mathbf{A} \cdot \mathbf{B})/(\mathbf{A} \cdot \mathbf{C}) \} \quad (2.19)$$

Considering $(\mathbf{B} \times \mathbf{C})$ as the third vector in the left-hand side of (2.13), the left-hand side of (2.19) is seen to be equal to $-\mathbf{A} \cdot \mathbf{B} \times (\mathbf{B} \times \mathbf{C})$. Therefore, applying (2.16), we obtain

$$\begin{aligned} \mathbf{B} \cdot \mathbf{A} \times (\mathbf{B} \times \mathbf{C}) &= -\mathbf{A} \cdot \mathbf{B} \times (\mathbf{B} \times \mathbf{C}) = -\mathbf{A} \cdot \{ (\mathbf{B} \cdot \mathbf{C}) \mathbf{B} - (\mathbf{B} \cdot \mathbf{B}) \mathbf{C} \} \\ &= (\mathbf{B} \cdot \mathbf{B})(\mathbf{A} \cdot \mathbf{C}) - (\mathbf{A} \cdot \mathbf{B})(\mathbf{B} \cdot \mathbf{C}) \end{aligned}$$

Comparing this with (2.19), we see that

$$k_1 = (\mathbf{A} \cdot \mathbf{C})$$

Finally, substitution into (2.18) gives the desired relation (2.14).

Let us now study the differentiations of vectors. There are two different kinds of differentiation, divergence and rotation, just as there are two multiplications, scalar product and vector product. The divergence ($\nabla \cdot \mathbf{A}$: del dot $\mathbf{A}$) is defined by

$$\nabla \cdot \mathbf{A} = \lim_{\Delta V \rightarrow 0} \frac{1}{\Delta V} \int_{\Delta S} \mathbf{A} \cdot \mathbf{n} dS \quad (2.20)$$

where the integral is over the closed surface of a small volume element ΔV and $\mathbf{n}$ is the outer normal unit vector, i.e., the unit vector normal to the surface and pointing outward. If $\mathbf{A}$ represents the velocity of incompressible liquid, since the integral gives the amount of the liquid emerging from the closed surface, $\nabla \cdot \mathbf{A}$ represents the amount of the liquid generated per unit volume. The divergence $\nabla \cdot \mathbf{A}$ is a scalar quantity which has no direction associated with it. The integration of $\nabla \cdot \mathbf{A}$ over a certain volume V is equivalent to the summation of $\nabla \cdot \mathbf{A} \Delta V$ over the same volume. Because of the definition, $\nabla \cdot \mathbf{A} \Delta V$ can be written in the form of a surface integral. Since the outflow through a surface from ΔV means the inflow through the same surface for the adjacent ΔV , when summing the surface integrals, the contributions from such common surfaces cancel each other. Thus, there remains only the contribution from the outermost surface which is not shared by the neighboring volume elements:

$$\int_V \nabla \cdot \mathbf{A} dv = \int_S \mathbf{A} \cdot \mathbf{n} dS \quad (2.21)$$

where S is the closed surface of the volume V and $\mathbf{n}$ is the outer normal unit vector as shown in Fig. 2.10. This is called Gauss's theorem.

Let us consider how to express the divergence in terms of the rectangular components A_x , A_y , and A_z of the vector $\mathbf{A}$. The shape of ΔV is irrelevant to the final result, however, to avoid unnecessary complication, let ΔV be a cube whose edges are parallel to the x , y , and z axes as shown in Fig. 2.11.

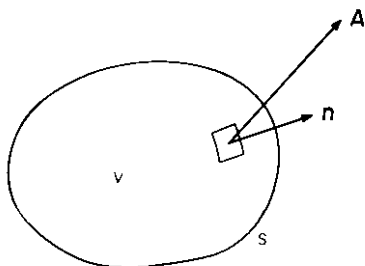


Fig. 2.10. Explanation of Gauss's theorem.

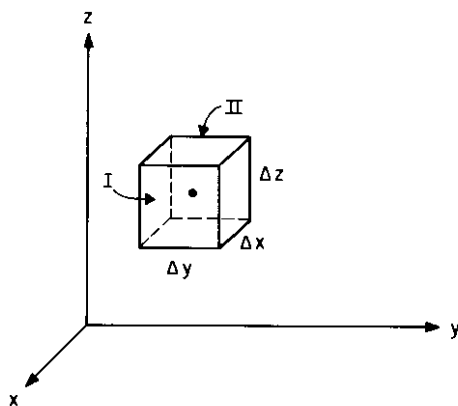


Fig. 2.11. Volume element for divergence calculation.

Let (x, y, z) be the center of the cube. The contribution to the surface integral from the two surfaces perpendicular to the x -axis is given by

$$\begin{aligned}
 & \int_I \mathbf{A} \cdot \mathbf{n} dS + \int_{II} \mathbf{A} \cdot \mathbf{n} dS \\
 &= A_x(x + \frac{1}{2} \Delta x, y, z) \Delta y \Delta z - A_x(x - \frac{1}{2} \Delta x, y, z) \Delta y \Delta z \\
 &= \{A_x(x, y, z) + \frac{1}{2} \Delta x (\partial A_x / \partial x)\} \Delta y \Delta z - \{A_x(x, y, z) - \frac{1}{2} \Delta x (\partial A_x / \partial x)\} \Delta y \Delta z \\
 &= (\partial A_x / \partial x) \Delta x \Delta y \Delta z
 \end{aligned}$$

Similarly, the contributions from other surfaces can be calculated, and when all are added, we have

$$\int_{\Delta S} \mathbf{A} \cdot \mathbf{n} dS = \left\{ \frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} \right\} \Delta x \Delta y \Delta z$$

Dividing by $\Delta V = \Delta x \Delta y \Delta z$ and taking the limit of $\Delta V \rightarrow 0$, we obtain

$$\nabla \cdot \mathbf{A} = \frac{\partial A_x}{\partial x} + \frac{\partial A_y}{\partial y} + \frac{\partial A_z}{\partial z} \quad (2.22)$$

Let ∇ be an operator defined by

$$\nabla = \mathbf{i}_x \frac{\partial}{\partial x} + \mathbf{i}_y \frac{\partial}{\partial y} + \mathbf{i}_z \frac{\partial}{\partial z} \quad (2.23)$$

Treating $\partial/\partial x$, $\partial/\partial y$, and $\partial/\partial z$ as if they were the rectangular components of a vector, when we calculate the scalar product of ∇ and $\mathbf{A}$, we get exactly the same expression as (2.22). For this reason, we indicate the divergence of $\mathbf{A}$ by $\nabla \cdot \mathbf{A}$.

Let us consider a special case in which $\mathbf{A}$ is a vector in a two dimensional space, i.e., $\mathbf{A}$ is parallel to the xy plane and is a function of x and y only (independent of z). Let V be a volume $S_0 \times 1$ where S_0 is the base area, and the height is unity as shown in Fig. 2.12. The integration of $\nabla \cdot \mathbf{A}$ over V is given by

$$\int_V \nabla \cdot \mathbf{A} dV = \int_{S_0} \nabla \cdot \mathbf{A} dS \times 1$$

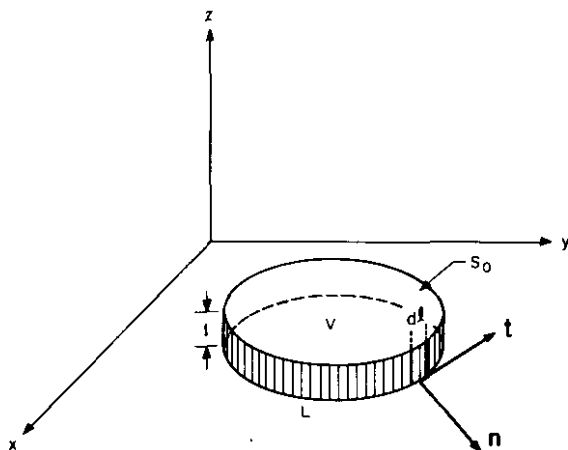


Fig. 2.12. Explanation for two-dimensional Gauss's theorem.

By Gauss's theorem, this is equal to the surface integral of $\mathbf{A} \cdot \mathbf{n}$ over the closed surface S of V . However, since by hypothesis, $\mathbf{A}$ is parallel to the xy plane, the contributions from the top and bottom surfaces vanish, and hence the surface integral over S is equal to the contribution from the cylindrical surface whose length is L :

$$\int_S \mathbf{A} \cdot \mathbf{n} dS = \int_L \mathbf{A} \cdot \mathbf{n} dl \times 1$$

From the above two equations and Gauss's theorem, we obtain

$$\int_{S_0} \nabla \cdot \mathbf{A} dS = \int_L \mathbf{A} \cdot \mathbf{n} dl \quad (2.24)$$

This is the two-dimensional Gauss's theorem. A comparison of (2.21) and (2.24) shows that the two dimensional case is obtainable simply by changing the volume and surface integrals in the ordinary Gauss's theorem to the surface and line integrals, respectively.

The other differentiation, the rotation of vector $\mathbf{A}$, is indicated by $\nabla \times \mathbf{A}$ (del cross $\mathbf{A}$). This is a vector quantity in contrast to the divergence which is a scalar quantity. The component of the rotation in the direction of an arbitrary unit vector $\mathbf{n}$ is given by

$$\mathbf{n} \cdot \nabla \times \mathbf{A} = \lim_{\Delta S \rightarrow 0} \frac{1}{\Delta S} \int_{\Delta C} \mathbf{A} \cdot d\mathbf{l} \quad (2.25)$$

where ΔS is a small area with $\mathbf{n}$ being normal to it, ΔC is the closed contour of ΔS and $d\mathbf{l}$ is the tangential vector of ΔC with a magnitude equal to the length of the small segment of ΔC , as illustrated in Fig. 2.13. The integration of $\mathbf{n} \cdot \nabla \times \mathbf{A}$ over a surface S is equivalent to the summation of a large number of quantities $\mathbf{n} \cdot \nabla \times \mathbf{A} \Delta S$ over the surface. Each can be converted

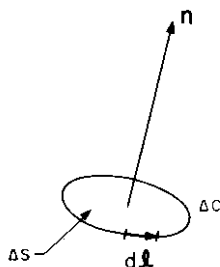


Fig. 2.13. Vector relation in the definition of rotation.

to the line integral along the closed contour ΔC . However, the contributions from the common borderline of the neighboring areas cancel each other since the directions of the integrals are opposite, as shown in Fig. 2.14. The remaining contribution comes from the outermost periphery of S only. Therefore, we obtain Stokes's theorem:

$$\int_S \nabla \times \mathbf{A} \cdot \mathbf{n} dS = \int_C \mathbf{A} \cdot d\mathbf{l} \quad (2.26)$$

where C is the closed contour of S . The direction of the contour integral is shown in Fig. 2.14.

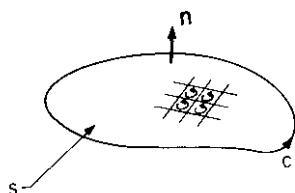


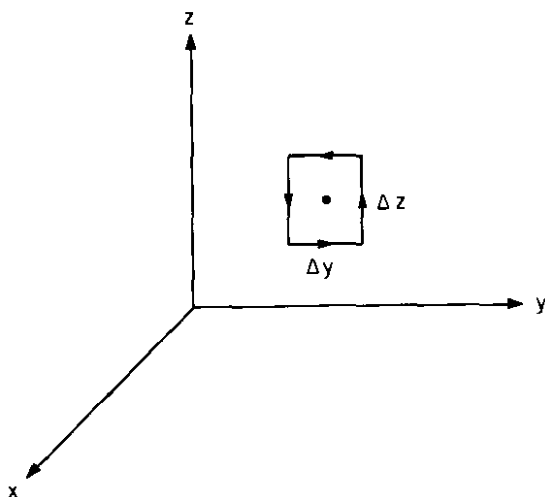
Fig. 2.14. Explanation of Stokes's theorem.

Let us consider how $\nabla \times \mathbf{A}$ should be expressed in terms of the rectangular components. If $\mathbf{n}$ is the unit vector in the x -direction, then, (2.25) gives the x -component of $\nabla \times \mathbf{A}$, namely $(\nabla \times \mathbf{A})_x$. The shape of ΔS is arbitrary as long as it is small and perpendicular to $\mathbf{n}$. However, in order to make the calculation simple, let us consider a rectangle with sides parallel to the y - and z -axes, as shown in Fig. 2.15. The center of the rectangle is given by (x, y, z) and the lengths of the sides are Δy and Δz , respectively. Then we have

$$\begin{aligned} (\nabla \times \mathbf{A})_x &= \lim (\Delta y \Delta z)^{-1} \int_{\Delta C} \mathbf{A} \cdot d\mathbf{l} \\ &= \lim (\Delta y \Delta z)^{-1} \{ A_z(x, y + \tfrac{1}{2} \Delta y, z) \Delta z - A_z(x, y - \tfrac{1}{2} \Delta y, z) \Delta z \\ &\quad + A_y(x, y, z - \tfrac{1}{2} \Delta z) \Delta y - A_y(x, y, z + \tfrac{1}{2} \Delta z) \Delta y \} \\ &= \lim (\Delta y \Delta z)^{-1} \{ (\partial A_z / \partial y) \Delta y \Delta z - (\partial A_y / \partial z) \Delta z \Delta y \} \\ &= (\partial A_z / \partial y) - (\partial A_y / \partial z) \end{aligned}$$

In a similar manner, other components can be obtained. Combining the results, we obtain

$$\nabla \times \mathbf{A} = \mathbf{i}_x \left(\frac{\partial A_z}{\partial y} - \frac{\partial A_y}{\partial z} \right) + \mathbf{i}_y \left(\frac{\partial A_x}{\partial z} - \frac{\partial A_z}{\partial x} \right) + \mathbf{i}_z \left(\frac{\partial A_y}{\partial x} - \frac{\partial A_x}{\partial y} \right) \quad (2.27)$$

Fig. 2.15. Integral contour for calculation of x -component of rotation.

The vector product of the operator ∇ defined in (2.23) and $\mathbf{A}$ gives the same expression; hence, the rotation of $\mathbf{A}$ is indicated by $\nabla \times \mathbf{A}$. Equation (2.27) can also be expressed in determinant form similar to (2.12):

$$\nabla \times \mathbf{A} = \begin{vmatrix} \mathbf{i}_x & \mathbf{i}_y & \mathbf{i}_z \\ \partial/\partial x & \partial/\partial y & \partial/\partial z \\ A_x & A_y & A_z \end{vmatrix} \quad (2.28)$$

Since $\nabla \times \mathbf{A}$ is a vector, we can consider the divergence of $\nabla \times \mathbf{A}$, $\nabla \cdot \nabla \times \mathbf{A}$. By definition,

$$\begin{aligned} \nabla \cdot \nabla \times \mathbf{A} &= \lim_{\Delta V \rightarrow 0} \frac{1}{\Delta V} \int_{\Delta S} \nabla \times \mathbf{A} \cdot \mathbf{n} \, dS \\ &= \lim_{\Delta V \rightarrow 0} \frac{1}{\Delta V} \int_{\Delta C} \mathbf{A} \cdot d\mathbf{l} \end{aligned}$$

where Stokes's theorem is used, ΔS is the closed surface of ΔV and ΔC is the closed contour of ΔS . However, since ΔS is a closed surface, the contour ΔC shrinks to a point and its length becomes zero. Thus, the contour integral is always equal to zero. This means that

$$\nabla \cdot \nabla \times \mathbf{A} = 0 \quad (2.29)$$

Since $\nabla \times \mathbf{A}$ is a vector, we can also consider the rotation of $\nabla \times \mathbf{A}$, $\nabla \times (\nabla \times \mathbf{A})$. However, in order to simplify the explanation, let us first

discuss a differentiation of scalar functions, which leads to a vector. A point in actual space can be specified by a vector drawn from the origin to the point. Therefore, let us identify each point by a vector $\mathbf{r}$. Let $\mathbf{r}_1$ and $\mathbf{r}_2$ be two points close to each other, and let φ be a single valued function whose values at $\mathbf{r}_1$ and $\mathbf{r}_2$ are indicated by φ_1 and φ_2 , respectively. In the limit of $\mathbf{r}_1 - \mathbf{r}_2 \equiv d\mathbf{r} \rightarrow 0$, let us express $d\varphi \equiv \varphi_1 - \varphi_2$ in the form of a scalar product of $d\mathbf{r}$ and some vector $\nabla\varphi$, i.e.,

$$d\varphi = \nabla\varphi \cdot d\mathbf{r} \quad (2.30)$$

The vector $\nabla\varphi$ defined through (2.30) is called the gradient of φ . Suppose $d\mathbf{r}$ lies on a constant φ surface, then $d\varphi$ is equal to zero and so is $\nabla\varphi \cdot d\mathbf{r}$ by definition. This means that $\nabla\varphi$ is perpendicular to $d\mathbf{r}$. Since the direction of $d\mathbf{r}$ is arbitrary as long as it lies on the surface, it follows that $\nabla\varphi$ is normal to the constant φ surface. When φ changes rapidly with distance normal to the constant φ surface, from the definition (2.30) the magnitude of $\nabla\varphi$ is large and vice versa. This is the reason why $\nabla\varphi$ is called the gradient of φ .

The integral of (2.30) along a closed path C must be equal to zero:

$$\int_C d\varphi = \int_C \nabla\varphi \cdot d\mathbf{r} = 0$$

This is because φ has the same value at both ends of the path which are really one and the same point. Since C is arbitrary, the definition of the rotation (2.25) shows that $\nabla \times \nabla\varphi$ is always equal to zero, i.e.,

$$\nabla \times \nabla\varphi = 0 \quad (2.31)$$

Since $\nabla\varphi$ is a vector, we can consider the divergence of $\nabla\varphi$, $\nabla \cdot \nabla\varphi$ which is sometimes written in the form $\nabla^2\varphi$.

Let us derive an expression for $\nabla\varphi$ in terms of its rectangular components. In the rectangular coordinate system, we have

$$\begin{aligned} d\mathbf{r} &= \mathbf{i}_x dx + \mathbf{i}_y dy + \mathbf{i}_z dz \\ d\varphi &= \frac{\partial\varphi}{\partial x} dx + \frac{\partial\varphi}{\partial y} dy + \frac{\partial\varphi}{\partial z} dz \end{aligned}$$

Substituting these expressions into (2.30), we obtain

$$\frac{\partial\varphi}{\partial x} dx + \frac{\partial\varphi}{\partial y} dy + \frac{\partial\varphi}{\partial z} dz = (\mathbf{i}_x \cdot \nabla\varphi) dx + (\mathbf{i}_y \cdot \nabla\varphi) dy + (\mathbf{i}_z \cdot \nabla\varphi) dz$$

Since dx, dy, dz are arbitrary infinitesimals, the above relation requires

$$\frac{\partial \varphi}{\partial x} = \mathbf{i}_x \cdot \nabla \varphi, \quad \frac{\partial \varphi}{\partial y} = \mathbf{i}_y \cdot \nabla \varphi, \quad \frac{\partial \varphi}{\partial z} = \mathbf{i}_z \cdot \nabla \varphi$$

This means that the x, y , and z components of $\nabla \varphi$ are given by $\partial \varphi / \partial x$, $\partial \varphi / \partial y$, and $\partial \varphi / \partial z$, respectively. Thus, we have

$$\nabla \varphi = \mathbf{i}_x \frac{\partial \varphi}{\partial x} + \mathbf{i}_y \frac{\partial \varphi}{\partial y} + \mathbf{i}_z \frac{\partial \varphi}{\partial z} \quad (2.32)$$

If φ is placed after the operator ∇ defined by (2.23), the same expression can be obtained, and it is for this reason that the gradient of φ is indicated by $\nabla \varphi$.

Since the divergence of a vector is obtainable by formally taking the scalar product of ∇ and the vector in the system of rectangular coordinates, the expression for $\nabla \cdot \nabla \varphi$ is given by

$$\nabla \cdot \nabla \varphi = \frac{\partial^2 \varphi}{\partial x^2} + \frac{\partial^2 \varphi}{\partial y^2} + \frac{\partial^2 \varphi}{\partial z^2} \quad (2.33)$$

which is simply the scalar product of ∇ and (2.32). The divergence of a vector $\mathbf{A}$ is a scalar function and we can, therefore, consider the gradient. This is expressed by $\nabla(\nabla \cdot \mathbf{A})$, where the parentheses indicate that the inside operation must be performed first.

We are now in a position to discuss the rotation of $\nabla \times \mathbf{A}$, which was previously postponed. Treating ∇ as if it were a vector defined by (2.23), the formal application of (2.14) to $\nabla \times (\nabla \times \mathbf{A})$ gives

$$\nabla \times (\nabla \times \mathbf{A}) = (\nabla \cdot \mathbf{A}) \nabla - (\nabla \cdot \nabla) \mathbf{A}$$

Since $\mathbf{B}$ in the first term on the right-hand side of (2.14) can be placed in front of $(\mathbf{A} \cdot \mathbf{C})$ without changing the relation, and since ∇ is a differential operator on $\mathbf{A}$ which is customarily placed in front of $\mathbf{A}$, let us move ∇ to come before $(\nabla \cdot \mathbf{A})$ in the first term on the right-hand side of the above equation. This gives

$$\nabla \times (\nabla \times \mathbf{A}) = \nabla(\nabla \cdot \mathbf{A}) - (\nabla \cdot \nabla) \mathbf{A} \quad (2.34)$$

The first term on the right-hand side is the gradient of the divergence $\mathbf{A}$. The second term is similar to the divergence of a gradient. Since $\mathbf{A}$ is a vector however, $(\nabla \cdot \nabla) \mathbf{A}$ has yet to be defined. If we use (2.34) to define $(\nabla \cdot \nabla) \mathbf{A}$ and also use ∇^2 to indicate $(\nabla \cdot \nabla)$, even when a vector follows it,

then we have

$$(\nabla \cdot \nabla) \mathbf{A} \equiv \nabla^2 \mathbf{A} = \nabla(\nabla \cdot \mathbf{A}) - \nabla \times (\nabla \times \mathbf{A}) \quad (2.35)$$

Since each term on the right-hand side is a vector, $\nabla^2 \mathbf{A}$ thus defined is a vector.

Let us derive an expression for $\nabla^2 \mathbf{A}$ in the rectangular coordinate system. In this system the divergence, gradient, and rotation are all obtained by the formal application of ∇ defined by (2.23) and since the definition of $\nabla^2 \mathbf{A}$ is obtained by the formal application of vector multiplication formula (2.14) to $\nabla \times (\nabla \times \mathbf{A})$, if we place $\mathbf{A}$ after the scalar multiplication of two ∇ 's as defined by (2.35) the expression for $\nabla^2 \mathbf{A}$ is obtained:

$$\begin{aligned} \nabla^2 \mathbf{A} &= \left(\mathbf{i}_x \frac{\partial}{\partial x} + \mathbf{i}_y \frac{\partial}{\partial y} + \mathbf{i}_z \frac{\partial}{\partial z} \right) \cdot \left(\mathbf{i}_x \frac{\partial}{\partial x} + \mathbf{i}_y \frac{\partial}{\partial y} + \mathbf{i}_z \frac{\partial}{\partial z} \right) \mathbf{A} \\ &= \frac{\partial^2 \mathbf{A}}{\partial x^2} + \frac{\partial^2 \mathbf{A}}{\partial y^2} + \frac{\partial^2 \mathbf{A}}{\partial z^2} \end{aligned} \quad (2.36)$$

In the rectangular coordinate system, $\nabla^2 \phi$ and $\nabla^2 \mathbf{A}$ have exactly the same form, as we can see from (2.33) and (2.36). This is not necessarily true, however, in other coordinate systems. When $\nabla^2 \mathbf{A}$ is defined through (2.35), it generally has a different differential form from that for $\nabla^2 \phi$ in the same coordinate system.

Let us next consider the differentiation of the product of functions. In ordinary differentiation, the derivative of a product of functions is given by the sum of terms in which only one of the functions is differentiated with the others remaining unchanged. For example, the derivative of fg is given by $f'g + fg'$ where the prime indicates the derivative. Since ∇ can be expressed in the form of (2.23), a similar operation is possible. Indicating the differential operators on $\mathbf{A}$ and on $\mathbf{B}$ by ∇_A and ∇_B , respectively, we have

$$\nabla \cdot (\mathbf{A} \nabla \cdot \mathbf{B}) = \nabla_A \cdot (\mathbf{A} \nabla \cdot \mathbf{B}) + \nabla_B \cdot (\mathbf{A} \nabla \cdot \mathbf{B}) = (\nabla \cdot \mathbf{A})(\nabla \cdot \mathbf{B}) + \mathbf{A} \cdot \nabla \nabla \cdot \mathbf{B} \quad (2.37)$$

Similarly, we have

$$\nabla \cdot (\mathbf{A} \times \nabla \times \mathbf{B}) = \nabla \times \mathbf{A} \cdot \nabla \times \mathbf{B} - \mathbf{A} \cdot \nabla \times \nabla \times \mathbf{B} \quad (2.38)$$

The negative sign in front of the second term on the right-hand side is the result of interchanging the order of ∇ and $\mathbf{A}$ during the derivation. The rectangular coordinate system was used to derive (2.37) and (2.38). Once both sides of an equation become vector expressions, however, the equation is valid, independent of the coordinate system utilized. If we take the volume

integrals of (2.37) and (2.38), by Gauss's theorem, the left-hand sides become the surface integrals. After the proper transposition of terms, we have

$$\int_V \mathbf{A} \cdot \nabla \nabla \cdot \mathbf{B} dV = - \int_V (\nabla \cdot \mathbf{A}) (\nabla \cdot \mathbf{B}) dV + \int_S (\mathbf{n} \cdot \mathbf{A}) (\nabla \cdot \mathbf{B}) dS \quad (2.39)$$

$$\int_V \mathbf{A} \cdot \nabla \times \nabla \times \mathbf{B} dV = \int_V \nabla \times \mathbf{A} \cdot \nabla \times \mathbf{B} dV - \int_S \mathbf{A} \times \nabla \times \mathbf{B} \cdot \mathbf{n} dS \quad (2.40)$$

These are formulas of integration by parts. When $\mathbf{A}$ and $\mathbf{B}$ are two-dimensional vectors, one can obtain the corresponding formulas by changing volume integrals to surface integrals and surface integrals to line integrals just as we did in the two-dimensional Gauss's theorem.

$$\int_S \mathbf{A} \cdot \nabla \nabla \cdot \mathbf{B} dS = - \int_S (\nabla \cdot \mathbf{A}) (\nabla \cdot \mathbf{B}) dS + \int_L (\mathbf{n} \cdot \mathbf{A}) (\nabla \cdot \mathbf{B}) dl \quad (2.41)$$

$$\int_S \mathbf{A} \cdot \nabla \times \nabla \times \mathbf{B} dS = \int_S \nabla \times \mathbf{A} \cdot \nabla \times \mathbf{B} dS - \int_L \mathbf{A} \times \nabla \times \mathbf{B} \cdot \mathbf{n} dl \quad (2.42)$$

The differentiation of a function of a function can also be discussed in a similar manner to the ordinary differentiation of a function of a function. As an example, let us consider ∇r where r is a scalar given by $(\mathbf{r} \cdot \mathbf{r})^{1/2}$. Since ∇r^2 is in the form of the differentiation of a function of a function, it is given by $2\mathbf{r} \cdot \nabla r$. On the other hand, from the differentiation of a product of functions, we have

$$\begin{aligned} \nabla r^2 &= \nabla(\mathbf{r} \cdot \mathbf{r}) = 2(\mathbf{r} \cdot \nabla) \mathbf{r} \\ &= 2 \{x(\partial/\partial x) + y(\partial/\partial y) + z(\partial/\partial z)\} (\mathbf{i}_x x + \mathbf{i}_y y + \mathbf{i}_z z) = 2\mathbf{r} \end{aligned}$$

Combining these two results, we obtain

$$\nabla r = (\mathbf{r}/r) \quad (2.43)$$

The same result is obtainable through the following manipulation.

$$\nabla r = \nabla(\mathbf{r} \cdot \mathbf{r})^{1/2} = \frac{1}{2}(\mathbf{r} \cdot \mathbf{r})^{-1/2} 2\mathbf{r} \cdot \nabla r = (\mathbf{r}/r)$$

where the relation $\mathbf{r} \cdot \nabla r = r$ calculated for (2.43) is used.

Let us apply (2.43) to the calculation of $\nabla^2(1/r)$. Since

$$\nabla(1/r) = -\nabla r/r^2 = -\mathbf{r}/r^3 \quad (2.44)$$

we have

$$\nabla \cdot \nabla \frac{1}{r} = -\nabla \cdot \frac{\mathbf{r}}{r^3} = -\frac{\nabla \cdot \mathbf{r}}{r^3} - \mathbf{r} \cdot \nabla \frac{1}{r^3} = -\frac{3}{r^3} + 3\mathbf{r} \cdot \frac{\nabla \mathbf{r}}{r^4} = -\frac{3}{r^3} + 3 \frac{\mathbf{r} \cdot \mathbf{r}}{r^5} = 0 \quad (2.45)$$

where use is made of the relation

$$\nabla \cdot \mathbf{r} = \{\mathbf{i}_x(\partial/\partial x) + \mathbf{i}_y(\partial/\partial y) + \mathbf{i}_z(\partial/\partial z)\} \cdot (\mathbf{i}_x x + \mathbf{i}_y y + \mathbf{i}_z z) = 3$$

Our next task is to derive Helmholtz's theorem which states: A differentiable but otherwise arbitrary vector function can always be expressed as the sum of the gradient of a scalar function and the rotation of a vector function. Before proceeding, let us first prove the following lemma: One of the solutions of Poisson's equation

$$\nabla^2 \varphi = -q \quad (2.46)$$

is given by

$$\varphi(\mathbf{r}) = \int \frac{q(\mathbf{r}')}{4\pi R} dv' \quad (2.47)$$

where R is defined by

$$R = \{(\mathbf{r} - \mathbf{r}') \cdot (\mathbf{r} - \mathbf{r}')\}^{1/2} \quad (2.48)$$

We shall indicate the differentiations with respect to $\mathbf{r}$ and $\mathbf{r}'$ by ∇ and ∇' , respectively. Let us apply ∇ to (2.47) and then integrate the result over a closed surface S . Interchanging the order of integrations on the right-hand side, we have

$$\int_S \nabla \varphi \cdot \mathbf{n} dS = \iiint_S \frac{q(\mathbf{r}')}{4\pi} \nabla \frac{1}{R} \cdot \mathbf{n} dS dv' \quad (2.49)$$

If $\mathbf{r}'$ is located outside the S , the surface integral on the right-hand side can be converted to a volume integral, i.e.,

$$\int_S \frac{q(\mathbf{r}')}{4\pi} \nabla \frac{1}{R} \cdot \mathbf{n} dS = \int_V \frac{q(\mathbf{r}')}{4\pi} \nabla \cdot \nabla \frac{1}{R} dv \quad (2.50)$$

where V is the volume enclosed by S . Since $\nabla \cdot \nabla(1/R)$ is equivalent to $\nabla \cdot \nabla(1/r)$ with origin at $\mathbf{r}'$, from (2.45) $\nabla \cdot \nabla(1/R)$ is equal to zero provided that R does not vanish. From this it follows that, when $\mathbf{r}'$ is outside S , the left-hand side of (2.50) is equal to zero. If $\mathbf{r}'$ should be inside S , R becomes zero at $\mathbf{r} = \mathbf{r}'$ and the integrand on the right-hand side of (2.49) loses the meaning at this point. To avoid this difficulty, let us consider a small sphere centered at $\mathbf{r}'$ with volume v_0 and surface S_0 . Using Gauss's theorem,

we have

$$\int_S \frac{q(\mathbf{r}')}{4\pi} \nabla \frac{1}{R} \cdot \mathbf{n} dS = \int_{V-v_0} \frac{q(\mathbf{r}')}{4\pi} \nabla \cdot \nabla \frac{1}{R} dV - \int_{S_0} \frac{q(\mathbf{r}')}{4\pi} \nabla \frac{1}{R} \cdot \mathbf{n} dS \quad (2.51)$$

where $\mathbf{n}$ is the outer normal unit vector from $V - v_0$. Note that $\mathbf{n}$ is directed toward the center of the sphere on S_0 . The first term on the right-hand side of (2.51) is equal to zero since $\nabla \cdot \nabla(1/R)$ vanishes everywhere in the three-dimensional region defined by $(V - v_0)$. The second term can be calculated as follows. Since

$$\mathbf{n} = -(\mathbf{r} - \mathbf{r}')/R$$

on S_0 , and since (2.44) clearly shows

$$\nabla(1/R) = -(\mathbf{r} - \mathbf{r}')/R^3$$

we have

$$\int_{S_0} \frac{q(\mathbf{r}')}{4\pi} \nabla \frac{1}{R} \cdot \mathbf{n} dS = \int_{S_0} \frac{q(\mathbf{r}')}{4\pi R^2} dS.$$

The integrand on the right-hand side is a constant on S_0 and the surface area of S_0 is $4\pi R^2$, hence the value of the integral is equal to $q(\mathbf{r}')$. From this and (2.51), it follows that when $\mathbf{r}'$ is inside S ,

$$\int_S \frac{q(\mathbf{r}')}{4\pi} \nabla \frac{1}{R} \cdot \mathbf{n} dS = -q(\mathbf{r}')$$

If $\mathbf{r}'$ is outside S , the surface integral is equal to zero as we showed before. Therefore, the contribution to the volume integral on the right-hand side of (2.49) comes from the elementary volume dv' inside S , i.e., from the volume V only. That is,

$$\int_S \nabla \varphi \cdot \mathbf{n} dS = - \int_V q(\mathbf{r}') dv'$$

By Gauss's theorem, the left-hand side becomes the volume integral of $\nabla \cdot \nabla \varphi$. The primes in the left-hand side can be omitted without changing the result. Therefore, we obtain

$$\int_V \{\nabla \cdot \nabla \varphi + q(\mathbf{r})\} dv = 0$$

Since V is arbitrary, it follows from this that (2.47) satisfies (2.46). This completes the proof of the lemma.

In order to get Helmholtz's theorem, we next apply the above result to the following vector version of Poisson's equation

$$\nabla^2 \mathbf{W} = -\mathbf{F} \quad (2.52)$$

Since in the rectangular coordinate system, (2.52) is satisfied if each component satisfies Poisson's equation (2.46),

$$\mathbf{W} = \int \frac{\mathbf{F}(\mathbf{r}')}{4\pi R} dv' \quad (2.53)$$

must be a solution of (2.52). On the other hand, from (2.35) and (2.52,) we have

$$\mathbf{F} = -\nabla^2 \mathbf{W} = -\nabla(\nabla \cdot \mathbf{W}) + \nabla \times (\nabla \times \mathbf{W}) \quad (2.54)$$

To obtain an explicit form of $\mathbf{F}$, let us calculate the divergence and rotation of (2.53) and substitute into the above equation. $-\nabla \cdot \mathbf{W}$ is given by

$$-\nabla \cdot \mathbf{W} = -\int \frac{\mathbf{F} \cdot \nabla}{4\pi} \frac{1}{R} dv' = \int \frac{\mathbf{F} \cdot \nabla'}{4\pi} \frac{1}{R} dv' \quad (2.55)$$

Using the relation

$$\nabla' \cdot \frac{\mathbf{F}}{R} = \frac{\nabla' \cdot \mathbf{F}}{R} + \mathbf{F} \cdot \nabla' \frac{1}{R}$$

we can rewrite the right-hand side of (2.55) and then apply Gauss's theorem. The result is

$$-\nabla \cdot \mathbf{W} = \int \frac{\mathbf{F} \cdot \mathbf{n}}{4\pi R} dS' - \int \frac{\nabla' \cdot \mathbf{F}}{4\pi R} dv' \quad (2.56)$$

On the other hand, $\nabla \times \mathbf{W}$ is given by

$$\nabla \times \mathbf{W} = -\int \frac{\mathbf{F}}{4\pi} \times \nabla \frac{1}{R} dv' = \int \frac{\mathbf{F}}{4\pi} \times \nabla' \frac{1}{R} dv'$$

Using the relation

$$\nabla' \times \frac{\mathbf{F}}{R} = \frac{\nabla' \times \mathbf{F}}{R} - \mathbf{F} \times \nabla' \frac{1}{R}$$

the above equation becomes

$$\nabla \times \mathbf{W} = \int \frac{\nabla' \times \mathbf{F}}{4\pi R} dv' - \int \nabla' \times \frac{\mathbf{F}}{4\pi R} dv' \quad (2.57)$$

If $\mathbf{K}$ is a constant vector, using Gauss's theorem, we have

$$\begin{aligned} \int \nabla' \cdot \mathbf{K} \times \frac{\mathbf{F}}{4\pi R} dv' &= -\mathbf{K} \cdot \int \nabla' \times \frac{\mathbf{F}}{4\pi R} dv' \\ &= \int \mathbf{K} \times \frac{\mathbf{F}}{4\pi R} \cdot \mathbf{n} dS' = \mathbf{K} \cdot \int \frac{\mathbf{F}}{4\pi R} \times \mathbf{n} dS' \end{aligned}$$

Since $\mathbf{K}$ is arbitrary, it follows from this that

$$-\int \nabla' \times \frac{\mathbf{F}}{4\pi R} dv' = \int \frac{\mathbf{F}}{4\pi R} \times \mathbf{n} dS'$$

Substituting this into the right-hand side of (2.57), we have

$$\nabla \times \mathbf{W} = \int \frac{\nabla' \times \mathbf{F}}{4\pi R} dv' + \int \frac{\mathbf{F} \times \mathbf{n}}{4\pi R} dS' \quad (2.58)$$

The primes on the right-hand sides of (2.56) and (2.58) can be omitted without changing these results. Substituting (2.56) and (2.58) into (2.54), we obtain the desired expression for Helmholtz's theorem,

$$\mathbf{F} = \nabla \left(\int \frac{\mathbf{F} \cdot \mathbf{n}}{4\pi R} dS - \int \frac{\nabla \cdot \mathbf{F}}{4\pi R} dv \right) + \nabla \times \left(\int \frac{\mathbf{F} \times \mathbf{n}}{4\pi R} dS + \int \frac{\nabla \times \mathbf{F}}{4\pi R} dv \right) \quad (2.59)$$

This shows that an arbitrary, but differentiable, vector function $\mathbf{F}$ can be expressed as the sum of the gradient of the scalar function in the parentheses of the first term and the rotation of the vector function in the second term. As is easily seen, $\nabla \cdot \mathbf{F} = 0$ (in volume V) alone is not sufficient for $\mathbf{F}$ to be expressible as the rotation of a vector function. On the other hand, $\nabla \cdot \mathbf{F} = 0$ (in volume V) and $\mathbf{F} \cdot \mathbf{n} = 0$ (on surface S) are sufficient to ensure $\mathbf{F}$ to be the rotation of a vector. Similarly, $\nabla \times \mathbf{F} = 0$ (in volume V) alone is not sufficient, but together with $\mathbf{F} \times \mathbf{n} = 0$ (on surface S), it guarantees that $\mathbf{F}$ is the gradient of a scalar function.

Before closing this section, let us derive the expressions for $\nabla \cdot \mathbf{A}$, $\nabla \varphi$, and $\nabla^2 \varphi$ in the cylindrical coordinate system, for later use. We apply the definition (2.20) for the divergence to the volume element ΔV given by $\Delta r \cdot r \Delta \theta \cdot \Delta z$ shown in Fig. 2.16. The surface integral is given by

$$\begin{aligned} &\{A_r(r + \Delta r) \cdot \Delta z \cdot (r + \Delta r) \Delta \theta - A_r(r) \Delta z \cdot r \Delta \theta\} \\ &\quad + \{A_\theta(\theta + \Delta \theta) \Delta r \cdot \Delta z - A_\theta(\theta) \Delta r \cdot \Delta z\} \\ &\quad + \{A_z(z + \Delta z) \Delta r \cdot r \Delta \theta - A_z(z) \Delta r \cdot r \Delta \theta\} \\ &= (\partial A_r / \partial r) \Delta r \Delta z \Delta \theta + (\partial A_\theta / \partial \theta) \Delta r \Delta z \Delta \theta + (\partial A_z / \partial z) \Delta r \Delta z r \Delta \theta \end{aligned}$$

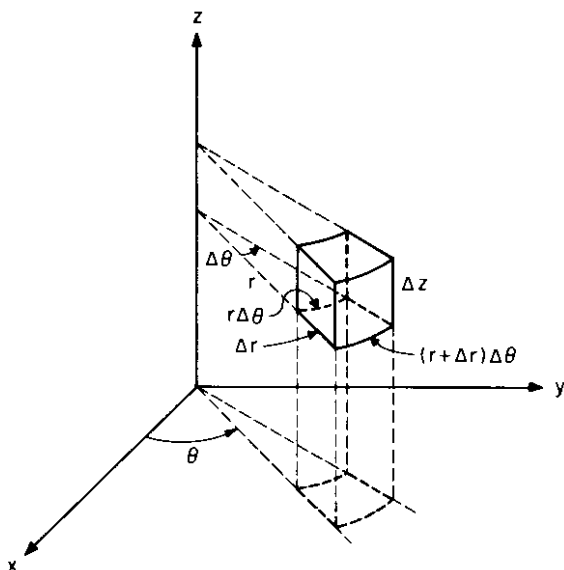


Fig. 2.16. Volume element ΔV in cylindrical coordinate system.

where the first, second, and third terms express the contribution from the r -constant, θ -constant, and z -constant surfaces, respectively, and use is made of the abbreviated forms of functions in which only that coordinate of primary concern is explicitly given in the parentheses. Dividing the above expression by ΔV , we obtain

$$\nabla \cdot \mathbf{A} = \frac{1}{r} \frac{\partial}{\partial r} (r A_r) + \frac{1}{r} \frac{\partial A_\theta}{\partial \theta} + \frac{\partial A_z}{\partial z} \quad (2.60)$$

In order to obtain the expression for $\nabla \phi$, let $\mathbf{i}_r$, $\mathbf{i}_\theta$, and $\mathbf{i}_z$ be the unit vectors in the r , θ , and z directions. The elementary displacement $d\mathbf{r}$ is given by

$$d\mathbf{r} = \mathbf{i}_r dr + \mathbf{i}_\theta r d\theta + \mathbf{i}_z dz$$

The increment $d\phi$ of a function ϕ due to the displacement $d\mathbf{r}$ is given by

$$\begin{aligned} d\phi &= \frac{\partial \phi}{\partial r} dr + \frac{\partial \phi}{\partial \theta} d\theta + \frac{\partial \phi}{\partial z} dz \\ &= \left(\mathbf{i}_r \frac{\partial \phi}{\partial r} + \mathbf{i}_\theta \frac{1}{r} \frac{\partial \phi}{\partial \theta} + \mathbf{i}_z \frac{\partial \phi}{\partial z} \right) \cdot (\mathbf{i}_r dr + \mathbf{i}_\theta r d\theta + \mathbf{i}_z dz) \\ &= \nabla \phi \cdot d\mathbf{r} \end{aligned}$$

It follows from this that

$$\nabla\varphi = \mathbf{i}_r \frac{\partial\varphi}{\partial r} + \mathbf{i}_\theta \frac{1}{r} \frac{\partial\varphi}{\partial\theta} + \mathbf{i}_z \frac{\partial\varphi}{\partial z} \quad (2.61)$$

Substituting this in place of $\mathbf{A}$ in (2.60), we obtain

$$\nabla \cdot \nabla\varphi = \frac{1}{r} \frac{\partial}{\partial r} \left(r \frac{\partial\varphi}{\partial r} \right) + \frac{1}{r^2} \frac{\partial^2\varphi}{\partial\theta^2} + \frac{\partial^2\varphi}{\partial z^2} \quad (2.62)$$

which is the expression for $\nabla^2\varphi$ in the cylindrical coordinate system.

2.2 Maxwell's Equations

If we take a coil having many turns with a voltmeter connected between its terminals and then move a magnet so that its flux threads the coil, the pointer of the voltmeter moves, indicating that some voltage is developed between the terminals. Since the voltage disappears when the magnet is stationary, the terminal voltage seems to be related to the time variation of the magnetic flux passing through the coil. By increasing the size of magnet, reversing the polarity, or by varying the speed with which the magnet approaches the coil, one is gradually convinced that the terminal voltage is proportional to the time derivative of the total flux crossing the coil, $\partial\Phi/\partial t$. Keeping the motion of the magnet relative to the coil the same, if the number of turns of the coil is increased, the terminal voltage increases proportionally. From this, it is reasonable to assume that a voltage appears in a coil with only one turn, which is proportional to the time derivative of the flux through the turn, i.e.,

$$V \propto \partial\Phi/\partial t$$

where $\propto$ indicates proportionality. On the right-hand side Φ can be expressed as the integral of the scalar product of magnetic flux density $\mathbf{B}$, i.e., magnetic induction, and the normal unit vector $\mathbf{n}$ over a surface S surrounded by the contour C which coincides with the coil as shown in Fig. 2.17. Also V is equal to the negative of the integral of electric field $\mathbf{E}$ from p to p' , however, the existence of V depends more on the coil surrounding the flux Φ than on the gap between p and p' . Since $\mathbf{E}$ must be small along the conductor, the limit of the integral from p to p' can be extended to include the entire closed contour C without changing the value appreciably. The above

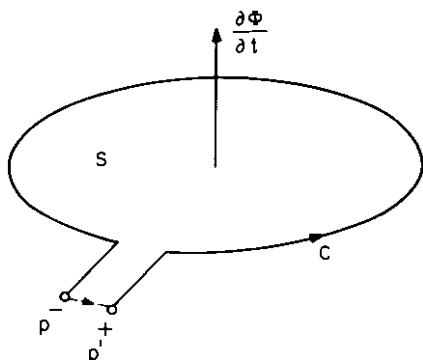


Fig. 2.17. Explanation of electromagnetic induction.

proportionality relation can then be rewritten in the form

$$-\int_C \mathbf{E} \cdot d\mathbf{l} \propto \frac{\partial}{\partial t} \int_S \mathbf{B} \cdot \mathbf{n} dS$$

Let us choose the unit of $\mathbf{B}$ so that the constant of proportionality becomes unity, i.e.,

$$\int_C \mathbf{E} \cdot d\mathbf{l} = -\frac{\partial}{\partial t} \int_S \mathbf{B} \cdot \mathbf{n} dS \quad (2.63)$$

where the direction of C is such that when a right-handed screw is turned in that direction, it will advance in the direction of $\mathbf{n}$.

Next, suppose a conducting wire carrying a heavy current passes vertically through the center of a sheet of paper with iron powder scattered on it. The iron powder will indicate that a magnetic field surrounds the current. A more elaborate experiment tells us that the magnetic field intensity is proportional to the magnitude of the current I and is inversely proportional to the distance from the conductor. Consider the contour integral of magnetic field $\mathbf{H}$ along a circle C_0 with the center at the conductor and radius r on a plane perpendicular to the conductor. Since the length of the contour is proportional to r and the intensity of $\mathbf{H}$ inversely proportional to r , the value of the integral becomes independent of r . However, since $\mathbf{H}$ is proportional to I , we have

$$\int_{C_0} \mathbf{H} \cdot d\mathbf{l} \propto I$$

The current can be expressed as the surface integral of the scalar product of

the current density $\mathbf{i}$ and the normal unit vector $\mathbf{n}$ over a surface S surrounded by C_0 . The above result then becomes similar to (2.63), for which the shape of C is arbitrary. Therefore, it may be natural to assume that

$$\int_C \mathbf{H} \cdot d\mathbf{l} = \int_S \mathbf{i} \cdot \mathbf{n} dS \quad (2.64)$$

where C is an arbitrary closed contour and S is a surface surrounded by C . The constant of proportionality is made unity by properly selecting the unit for $\mathbf{H}$.

Let us first check the validity of (2.64) for the special case of the vertical conducting wire discussed above. If we use the cylindrical coordinate system with the z -axis coinciding with the conductor, then since

$$\mathbf{H} \propto \mathbf{i}_\theta (I/r)$$

and

$$d\mathbf{l} = \mathbf{i}_r dr + \mathbf{i}_\theta r d\theta + \mathbf{i}_z dz$$

the contour integral of $\mathbf{H}$ becomes

$$\int_C \mathbf{H} \cdot d\mathbf{l} \propto \int_C I d\theta$$

which is equal to $2\pi I$ when C encloses the conductor and is equal to zero, otherwise. If 2π is included in the constant of proportionality, this result confirms (2.64) for this special case. However, if we assume that (2.64) is always correct, we encounter some difficulties, especially when a capacitor is inserted in the conducting path as shown in Fig. 2.18. A magnetic field

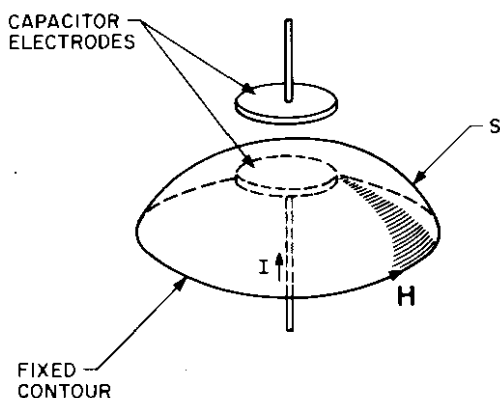


Fig. 2.18. Explanation of displacement current.

surrounds the conductor as long as current is flowing through it, and if S is chosen so as to cut the vertical conducting wire in Fig. 2.18, there is no problem. However, if S passes between the electrodes of the capacitor, as shown in Fig. 2.18, the current density $\mathbf{i}$ is zero everywhere on S , and hence (2.64) leads to a contradiction that the line integral of $\mathbf{H}$ along the same contour C becomes zero or finite, depending on how S is chosen. This necessitates some modification of (2.64). When current is flowing into a capacitor, it is gradually charged and the electric field between its electrodes increases with time. Therefore, $\partial\mathbf{D}/\partial t$ exists in place of $\mathbf{i}$, where $\mathbf{D} = \epsilon\mathbf{E}$ is the electric displacement. Let us call $\partial\mathbf{D}/\partial t$ the displacement current (although nothing is displaced in vacuum) and add this to the current in the right-hand side of (2.64) to take care of the above difficulty:

$$\int_C \mathbf{H} \cdot d\mathbf{l} = \int_S \left(\mathbf{i} + \frac{\partial\mathbf{D}}{\partial t} \right) \cdot \mathbf{n} dS \quad (2.65)$$

where the magnitude of ϵ is so chosen that the constant of proportionality in front of $\partial\mathbf{D}/\partial t$ becomes unity.

We have now obtained (2.63) and (2.65) as the most plausible equations expressing electromagnetic field relations. Although the left-hand sides of the two equations are contour integrals whereas the right-hand sides are surface integrals, we can convert the contour integrals to surface integrals by Stokes' theorem. After transpositions, we have

$$\int_S \left(\nabla \times \mathbf{E} + \frac{\partial\mathbf{B}}{\partial t} \right) \cdot \mathbf{n} dS = 0 \quad (2.66)$$

$$\int_S \left\{ \nabla \times \mathbf{H} - \left(\mathbf{i} + \frac{\partial\mathbf{D}}{\partial t} \right) \right\} \cdot \mathbf{n} dS = 0 \quad (2.67)$$

Since S is arbitrary, the integrands have to be zero:

$$\nabla \times \mathbf{E} = - \frac{\partial\mathbf{B}}{\partial t} \quad (2.68)$$

$$\nabla \times \mathbf{H} = \mathbf{i} + \frac{\partial\mathbf{D}}{\partial t} \quad (2.69)$$

Since there is always electric charge q at the terminal point of $\mathbf{D}$ which is considered as the source of $\mathbf{D}$, from the definition of divergence (2.20), we have

$$\nabla \cdot \mathbf{D} = q \quad (2.70)$$

where the unit of charge is again properly chosen so that the constant of proportionality is unity. On the other hand, positive and negative magnetic charges are always found in pairs, and in any finite volume there is no net magnetic charge. This experimental fact is expressed in the form

$$\mathbf{V} \cdot \mathbf{B} = 0 \quad (2.71)$$

The four equations (2.68), (2.69), (2.70), and (2.71) are called Maxwell's equations.

Taking the divergence of (2.69) and noting that the order of differentiations with respect to time and position can be interchanged, we obtain

$$\mathbf{V} \cdot \mathbf{i} + \frac{\partial}{\partial t} \mathbf{V} \cdot \mathbf{D} = 0$$

since the divergence of rotation is always equal to zero. The substitution of (2.70) into the equation gives

$$\mathbf{V} \cdot \mathbf{i} = - \frac{\partial q}{\partial t} \quad (2.72)$$

This expresses the conservation of charge since by Gauss's theorem it means that the decrease of q in a certain volume V is due to the outflow of q in the form of current through the surface S enclosing V . Similarly, taking the divergence of (2.68), it can be shown that $\mathbf{V} \cdot \mathbf{B}$ is a constant independent of time. This constant is chosen to be zero in (2.71) to conform with experiments.

As we have already discussed, the choice of the units for electromagnetic quantities is made in such a way that the constants of proportionality become unity in Maxwell's equations. The unit of each quantity appearing in the equations is listed below:

E:	[volt]/[meter]	[V/m]
H:	[ampere-turn]/[meter]	[AT/m]
B:	[weber]/[meter] ²	[Wb/m ²]
D:	[coulomb]/[meter] ²	[C/m ²]
i:	[ampere]/[meter] ²	[A/m ²]
q:	[coulomb]/[meter] ³	[C/m ³]

Maxwell derived his equations through more or less the same argument described above after being stimulated by the experimental results of Faraday and others. Mathematically, the equations do not contradict each other, but

this does not mean that they are proved. In fact, there is no proof that electromagnetic phenomena should obey Maxwell's equations which must, therefore, be considered as bold postulates. The only support for the validity of Maxwell's equations is the fact that no macroscopic experiment has been conducted successfully which disproves them. Although there is no proof, let us accept their validity and consider how various phenomena can be explained or what kind of phenomena should be expected as natural consequences of Maxwell's equations. For instance, one argument goes as follows: Equation (2.69) gives (2.67) which in turn becomes (2.65), and it follows from this and a symmetry argument that the intensity of $\mathbf{H}$ around a straight conductor carrying a current I is proportional to I and inversely proportional to the distance from the conductor. Although this explanation of Maxwell's equations might sound strange to some of us, it should be remembered that physics is generally based upon similar foundations. Let us consider Newtonian mechanics. From experiments, it was inferred that force $\mathbf{f}$ was proportional to mass m times acceleration, i.e.,

$$\mathbf{f} = m \frac{d^2 \mathbf{r}}{dt^2} \quad (2.73)$$

where $\mathbf{r}$ indicates the position. By assuming that this equation always holds, a number of useful theorems were deduced from it. There was no proof for (2.73), in fact, it was later discovered that when the velocity of an object approached light velocity, (2.73) was no longer adequate and this led to the theory of relativity. In spite of these limitations, (2.73) is a very valuable result and worth studying together with the theorems deduced from it. Similarly, some macroscopic phenomena might be discovered in the future which disprove the universal validity of Maxwell's equations. Then, the equations will have to be revised accordingly. Nevertheless, the present form of Maxwell's equations will remain worth studying which is one of the objectives of this book.

In ordinary media, $\mathbf{D}$ and $\mathbf{B}$ are approximately proportional to $\mathbf{E}$ and $\mathbf{H}$, respectively. The constants of proportionality are generally indicated by ϵ and μ :

$$\mathbf{D} = \epsilon \mathbf{E} \quad (2.74)$$

$$\mathbf{B} = \mu \mathbf{H} \quad (2.75)$$

Since the units of $\mathbf{D}$ and $\mathbf{B}$ have been chosen so as to make the constants of proportionality in Maxwell's equations unity, ϵ and μ for a vacuum are not

unity but have the following values:

$$\begin{aligned}\epsilon_0 &= 8.854 \times 10^{-12} \quad [\text{farad/meter}] \\ \mu_0 &= 1.257 \times 10^{-6} \quad [\text{henry/meter}]\end{aligned}$$

where the subscript 0 is used to indicate the values for vacuum. In other media, ϵ and μ are different from ϵ_0 and μ_0 . Therefore, we write as

$$\epsilon = \epsilon_r \epsilon_0$$

$$\mu = \mu_r \mu_0$$

and call ϵ_r the relative dielectric constant and μ_r the relative permeability of the material. The current density is also approximately proportional to $\mathbf{E}$ in ordinary materials, i.e.,

$$\mathbf{i} = \sigma \mathbf{E} \quad (2.76)$$

The constant of proportionality σ is called the conductivity of the material.

Equations (2.74)–(2.76) are approximations. However, if we consider an idealized model for which these relations hold, then (2.68) and (2.69) become

$$\nabla \times \mathbf{E} = -\mu \frac{\partial \mathbf{H}}{\partial t} \quad (2.77)$$

$$\nabla \times \mathbf{H} = \sigma \mathbf{E} + \epsilon \frac{\partial \mathbf{E}}{\partial t} \quad (2.78)$$

These equations are linear and the coefficients are real, therefore the same arguments we used in the discussion of transmission line theory can be applied here, and the same time factor $e^{j\omega t}$ can be used for the electromagnetic field. The differentiation with respect to time becomes equivalent to the multiplication by $j\omega$. Thus, (2.77) and (2.78) become

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H} \quad (2.79)$$

$$\nabla \times \mathbf{H} = \sigma\mathbf{E} + j\omega\epsilon\mathbf{E} \quad (2.80)$$

These are Maxwell's equations in the present case. Equation (2.71) can be derived from (2.79), and (2.70) from the conservation of charge and (2.80).

Let us turn our attention to energy. Substituting (2.77) and (2.78) into the identity

$$\nabla \cdot \mathbf{E} \times \mathbf{H} = \mathbf{H} \cdot \nabla \times \mathbf{E} - \mathbf{E} \cdot \nabla \times \mathbf{H}$$

and integrating the result over a volume V , we obtain

$$\int_S \mathbf{E} \times \mathbf{H} \cdot \mathbf{n} \, dS = - \int_V \sigma \mathbf{E}^2 \, dv - \frac{\partial}{\partial t} \int_V \frac{1}{2} (\epsilon \mathbf{E}^2 + \mu \mathbf{H}^2) \, dv \quad (2.81)$$

where the left-hand side is converted into the surface integral by Gauss's theorem. The surface enclosing V is S , and $\mathbf{E}^2$ is an abbreviation for $\mathbf{E} \cdot \mathbf{E}$. The vector $\mathbf{E} \times \mathbf{H}$ appearing in the left integral is called Poynting vector.

Let us assume that there are electrodes at the top and bottom surfaces of a small cylindrical volume ΔV whose axis is parallel to $\mathbf{E}$, as shown in Fig. 2.19, and consider the power consumption in ΔV . The voltage between

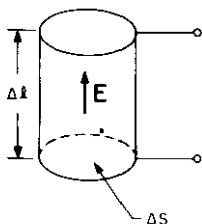


Fig. 2.19. Explanation of power consumption in conductive medium.

the electrodes is given by $|\mathbf{E}| \, \Delta l$ and the current by $\sigma |\mathbf{E}| \, \Delta S$. The ohmic loss is given by voltage times current which is equal to $\sigma \mathbf{E}^2 \, \Delta S \, \Delta l = \sigma \mathbf{E}^2 \, \Delta V$. Eliminating the electrodes, the same amount of power is considered being consumed since the electric field in ΔV stays the same. Therefore, $\int_V \sigma \mathbf{E}^2 \, dv$ should give the total power consumption in V . Since this term is related to energy, the other terms in (2.81) must also be related to energy. Although there is no particular reason for the following approach, for convenience, let us argue as follows: The integral of the Poynting vector over S , $\int_S \mathbf{E} \times \mathbf{H} \cdot \mathbf{n} \, dS$, represents the power flowing out through S , $\int_V \frac{1}{2} \epsilon \mathbf{E}^2 \, dv$ is the electric energy stored in V and $\int_V \frac{1}{2} \mu \mathbf{H}^2 \, dv$ is the magnetic energy stored in V . Then, (2.81) expresses the conservation of energy since the sum of the power consumed in V and the outflow of power through S becomes equal to the rate of decrease of the total stored energy.

Let us next consider the case in which the electric and magnetic fields have the time factor $e^{j\omega t}$, Maxwell's equations are then given by (2.79) and (2.80). The average power consumption in V can be calculated by multiplying $\mathbf{E}$ by $\sqrt{2} \, e^{j\omega t}$ and taking the real part for the interpretation of $\mathbf{E}$ in the

equations. Therefore, for $\sigma \mathbf{E}^2$, we have

$$\begin{aligned}\sigma \{\operatorname{Re}(\sqrt{2} \mathbf{E} e^{j\omega t})\}^2 &= 2\sigma \{\cos \omega t (\operatorname{Re} \mathbf{E}) - \sin \omega t (\operatorname{Im} \mathbf{E})\}^2 \\ &= 2\sigma \{\cos^2 \omega t (\operatorname{Re} \mathbf{E})^2 - 2 \cos \omega t \sin \omega t (\operatorname{Re} \mathbf{E}) \cdot (\operatorname{Im} \mathbf{E}) + \sin^2 \omega t (\operatorname{Im} \mathbf{E})^2\}\end{aligned}$$

We integrate this over one period T and divide the result by T to take the time-average. The average of $\cos^2 \omega t$ over a period is seen to be one-half from Fig. 2.20. Similarly, the average value of $\sin^2 \omega t$ is given by $\frac{1}{2}$, while the average of $\cos \omega t \sin \omega t$ is zero. Therefore, the average power consumption is given by

$$\begin{aligned}\int_V T^{-1} \int_0^T \sigma [\operatorname{Re}(\sqrt{2} \mathbf{E} e^{j\omega t})]^2 dt dv &= \int_V \sigma \{(\operatorname{Re} \mathbf{E})^2 + (\operatorname{Im} \mathbf{E})^2\} dv \\ &= \int_V \sigma \mathbf{E} \cdot \mathbf{E}^* dv\end{aligned}$$

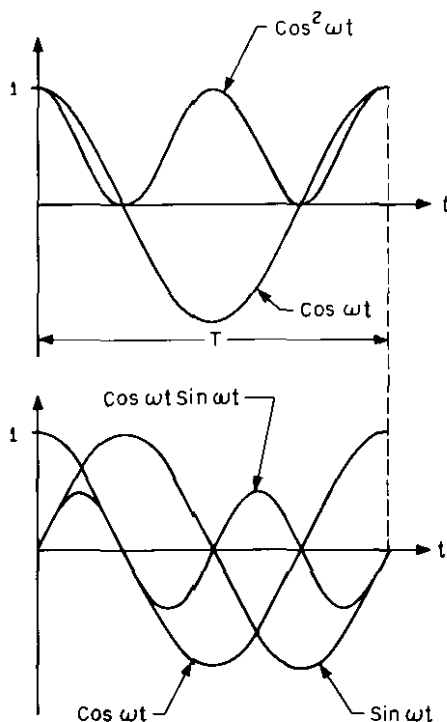


Fig. 2.20. Graphs of $\cos^2 \omega t$ and $\cos \omega t \sin \omega t$.

A similar discussion shows that the average electric and magnetic energies stored in V over a period are given by $\int_V \frac{1}{2} \epsilon \mathbf{E} \cdot \mathbf{E}^* dv$ and $\int_V \frac{1}{2} \mu \mathbf{H} \cdot \mathbf{H}^* dv$, respectively.

Now, subtracting $\mathbf{H}^* \cdot (2.79)$ from $\mathbf{E} \cdot \{\text{the complex conjugate of (2.80)}\}$ and integrating the result over V , we have

$$\begin{aligned} \int_V (\mathbf{E} \cdot \nabla \times \mathbf{H}^* - \mathbf{H}^* \cdot \nabla \times \mathbf{E}) dv \\ = \int_V \sigma \mathbf{E} \cdot \mathbf{E}^* dv + j\omega \int_V (\mu \mathbf{H} \cdot \mathbf{H}^* - \epsilon \mathbf{E} \cdot \mathbf{E}^*) dv \end{aligned} \quad (2.82)$$

Using the identity

$$\nabla \cdot \mathbf{E} \times \mathbf{H}^* = \mathbf{H}^* \cdot \nabla \times \mathbf{E} - \mathbf{E} \cdot \nabla \times \mathbf{H}^*$$

the left-hand side of (2.82) can be converted to a surface integral by Gauss's theorem. The result is

$$\int_S (\mathbf{E} \times \mathbf{H}^*) \cdot (-\mathbf{n}) dS = \int_V \sigma \mathbf{E} \cdot \mathbf{E}^* dv + j\omega \int_V (\mu \mathbf{H} \cdot \mathbf{H}^* - \epsilon \mathbf{E} \cdot \mathbf{E}^*) dv$$

$\mathbf{E} \times \mathbf{H}^*$ in the left integral is called the complex Poynting vector. The first integral on the right-hand side is the average power consumption while the second integral is twice the difference between the average magnetic and electric energy stored in V . Since the real and imaginary parts can be equated separately, we have

$$\text{Re} \int_S \mathbf{E} \times \mathbf{H}^* \cdot (-\mathbf{n}) dS = (\text{average power consumption in } V) \quad (2.83)$$

$$\begin{aligned} \text{Im} \int_S \mathbf{E} \times \mathbf{H}^* \cdot (-\mathbf{n}) dS = 2\omega (\text{average magnetic stored energy} \\ - \text{average electric stored energy}) \end{aligned} \quad (2.84)$$

The power consumed in V may be considered as coming through S , and if we assume that the incoming power per unit area is given by $\text{Re}(\mathbf{E} \times \mathbf{H}^*) \cdot (-\mathbf{n})$ at each point on S , the total power coming through S becomes exactly equal to the power consumption in V , as shown in (2.83). For this reason, let us consider that the real part of the complex Poynting vector represents the power flow density. However, this interpretation of Poynting vector is rather arbitrary. We could consider that the power flow per unit area was given by $\{\text{Re}(\mathbf{E} \times \mathbf{H}^*) + \nabla \times \mathbf{X}\} \cdot (-\mathbf{n})$ where $\mathbf{X}$ was an arbitrary single valued vector function of position. By Stokes's theorem, the integra-

tion of $\nabla \times \mathbf{X} \cdot (-\mathbf{n})$ over a closed surface always vanishes, and $\mathbf{X}$ does not appear in the final result.

2.3 Plane Waves

In this section, we shall consider electromagnetic waves which depend upon a single straight-space coordinate, i.e., waves whose amplitude and phase are constant over each plane perpendicular to the coordinate. Our first problem is to determine whether or not such electromagnetic fields can exist. To do so, we assume that $\mathbf{E}$ and $\mathbf{H}$ are independent of x and y in the rectangular coordinate system and then seek solutions of Maxwell's equation. Since $\mathbf{E}$ and $\mathbf{H}$ are independent of x and y , the derivatives $\partial/\partial x$ and $\partial/\partial y$ must be zero, and when the field vectors are decomposed into their components, Maxwell's equations (2.79) and (2.80) take the form

$$-(\partial H_y / \partial z) = (\sigma + j\omega\epsilon) E_x \quad (2.85)$$

$$(\partial H_x / \partial z) = (\sigma + j\omega\epsilon) E_y \quad (2.86)$$

$$0 = (\sigma + j\omega\epsilon) E_z \quad (2.87)$$

$$-(\partial E_y / \partial z) = -j\omega\mu H_x \quad (2.88)$$

$$(\partial E_x / \partial z) = -j\omega\mu H_y \quad (2.89)$$

$$0 = -j\omega\mu H_z \quad (2.90)$$

From (2.87) and (2.90), it is obvious that E_z and H_z become zero. Next, eliminating H_y from (2.85) and (2.89), we have

$$\frac{d^2 E_x}{dz^2} + (\omega^2 \epsilon \mu - j\omega \mu \sigma) E_x = 0 \quad (2.91)$$

where the symbol for ordinary differentiation is used instead of the one for partial differentiation since E_x is a function of z only. If E_x is obtained from this equation under appropriate boundary conditions, H_y is then determined from (2.89), E_y and H_x however, remain undetermined. If we eliminate H_x from (2.86) and (2.88), we obtain an equation for E_y identical to (2.91). This equation and appropriate boundary conditions determine E_y and H_x is then obtained from (2.88). Thus, E_x and H_y constitute one pair of related vectors, and E_y and H_x another. These pairs are generally independent of each other, except for possible interaction at the boundary. Since they are similar, let us first concentrate on the first pair E_x and H_y . Equation (2.91)

is an ordinary differential equation with constant coefficients. Therefore, if we assume that $E_x = Ae^{-\gamma z}$, just as we did in (1.7), and substitute it into (2.91), we have

$$\gamma^2 = -(\omega^2 \epsilon \mu - j\omega \mu \sigma) \quad (2.92)$$

For the two solutions of (2.92), let us use the expression

$$\gamma = \pm (\alpha + j\beta) \quad (2.93)$$

where α and β are real, and assume that β is positive. Since γ^2 , given by (2.92), is in the second quadrant of the complex plane, the γ 's are in the

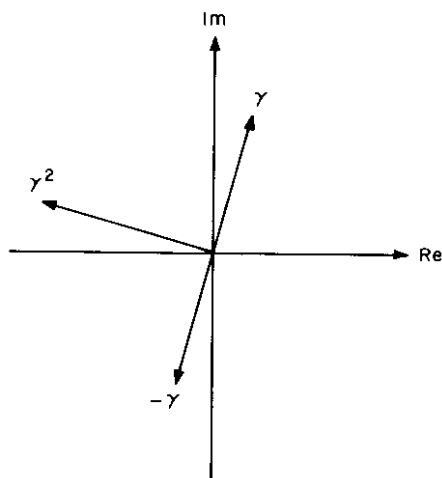


Fig. 2.21. Relation between γ^2 and γ .

first and third quadrants, as shown in Fig. 2.21. From this we see that α is also positive since β is assumed to be positive. From (2.92), we can write

$$\alpha + j\beta = j(\omega^2 \epsilon \mu - j\omega \mu \sigma)^{1/2} \quad (2.94)$$

where the real part of the square root is taken as positive. To eliminate unnecessary confusion in future, let us now agree that the real part of $()^{1/2}$ is always positive. When the other root is required we shall write it as $-()^{1/2}$. Note that the sign in (1.8) complies with this rule.

From the above discussion, the most general expression for E_x is given by

$$E_x = Ae^{-(\alpha + j\beta)z} + Be^{(\alpha + j\beta)z} \quad (2.95)$$

Substituting (2.95) into (2.89), we obtain

$$H_y = Z_0^{-1} \{Ae^{-(\alpha + j\beta)z} - Be^{(\alpha + j\beta)z}\} \quad (2.96)$$

where

$$Z_0 = [\mu/\{\epsilon - j(\sigma/\omega)\}]^{1/2} \quad (2.97)$$

Substituting (2.95) and (2.96) into (2.85), we see that (2.85) is also satisfied. Thus, we conclude that an electromagnetic field does exist which is independent of the x - and y -coordinates and satisfies Maxwell's equations.

Since it is the ratio of electric to magnetic fields Z_0 has the dimension of impedance; electric field times distance has the dimension of voltage, and magnetic field times distance has that of current. The quantity Z_0 is determined by the property of the medium only and is referred to as its characteristic impedance with a positive real part in accordance with the agreement on the sign of square root. Note that we used Z_0 in Section 1.1 with a different meaning. From a figure similar to Fig. 2.21, it can be shown that the imaginary part of Z_0 is also positive, in other words, Z_0 is inductive.

The constants A and B in (2.95) and (2.96) are determined by boundary conditions. First, however, let us assume $B = 0$ and concentrate on the terms with A only. The term $e^{-j\beta z}$ represents a wave traveling in the positive z direction, and $e^{-\alpha z}$ represents the exponential decay with distance. Since Z_0^{-1} has a negative imaginary part, let us write it in the form $|Z_0|^{-1} e^{-j\varphi}$. To interpret the magnetic field expression, we multiply it by $\sqrt{2} e^{j\omega t}$ and take the real part of the result:

$$\text{Re} \{ \sqrt{2} |Z_0|^{-1} e^{-j\varphi} A e^{-(\alpha + j\beta)z} e^{j\omega t} \} = \sqrt{2} |Z_0|^{-1} e^{-\alpha z} A \cos(\omega t - \beta z - \varphi)$$

Thus, the phase of magnetic field lags by φ behind that of electric field as depicted in Fig. 2.22. Similarly, the terms with B are found to represent an

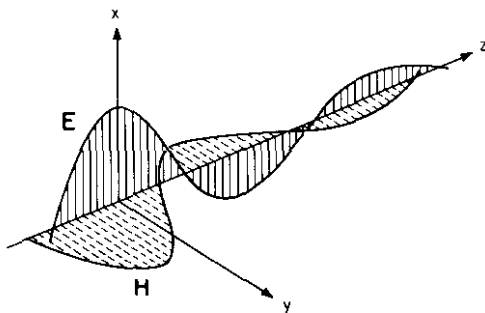


Fig. 2.22. Electromagnetic field of a plane wave.

electromagnetic wave traveling in the negative z direction with decreasing amplitude as it propagates. Since at each instant, the electric and magnetic fields are both constant over each plane perpendicular to the z axis, the above electromagnetic waves are called plane waves.

When σ is equal to zero, no attenuation takes place. Furthermore, the phase lag of the magnetic field disappears. In this case, the propagation constant and the characteristic impedance are given by

$$\alpha + j\beta = 0 + j\omega(\epsilon\mu)^{1/2} \quad (2.98)$$

$$Z_0 = (\mu/\epsilon)^{1/2} \quad (2.99)$$

respectively. The corresponding wavelength is

$$\lambda = 2\pi/\beta = 2\pi/\omega(\epsilon\mu)^{1/2}$$

On the other hand, when σ is large, ϵ can be neglected compared to σ/ω . Thus, we have

$$\alpha + j\beta \simeq (j\omega\mu\sigma)^{1/2} = (1 + j)(\omega\mu\sigma/2)^{1/2} \quad (2.100)$$

$$Z_0 \simeq \{\mu/(-j\sigma/\omega)\}^{1/2} = (1 + j)(\omega\mu/2\sigma)^{1/2} \quad (2.101)$$

This means that the phase of magnetic field lags 45° behind that of electric field and since the magnitude of Z_0 is small, the magnetic field predominates over the electric field, compared to the free space case. The amplitudes of both electric and magnetic fields decrease by a factor of $1/e$ when they travel a distance $1/\alpha = (2/\omega\mu\sigma)^{1/2}$ called the skin depth of the conductor.

Let us consider the transmission power per unit area. Since the electric and magnetic fields are orthogonal to each other and both are perpendicular to the direction of propagation, the Poynting vector is in the z direction, and its value is given by $E_x H_y^*$. Therefore, the transmission power is $\text{Re}\{E_x H_y^*\}$ per unit area.

Let a and b be defined by

$$a = \frac{1}{2} |\text{Re } Z_0|^{-1/2} (E_x + Z_0 H_y), \quad b = \frac{1}{2} |\text{Re } Z_0|^{-1/2} (E_x - Z_0^* H_y) \quad (2.102)$$

Then, a calculation similar to the one leading to (1.53) gives

$$|a|^2 - |b|^2 = \text{Re}\{E_x H_y^*\} \quad (2.103)$$

This means that $|a|^2 - |b|^2$ is equal to the transmission power per unit area. Substituting (2.95) and (2.96) into (2.102), we have

$$a = |\text{Re } Z_0|^{-1/2} A e^{-(\alpha + j\beta)z} \quad (2.104)$$

$$b = |\operatorname{Re} Z_0|^{-1/2} \left\{ \frac{\operatorname{Im} Z_0}{Z_0} A e^{-(\alpha + j\beta)z} + \frac{\operatorname{Re} Z_0}{Z_0} B e^{(\alpha + j\beta)z} \right\} \quad (2.105)$$

When $\sigma = 0$, Z_0 becomes real, and the first term in the parentheses on the right-hand side of (2.105) disappears. In this case, a and b can be interpreted as the waves traveling in the positive z and negative z directions, respectively, and the square of the magnitude of each wave gives the power transmitted in the direction of propagation. On the other hand, when $\sigma \neq 0$, b can no longer be interpreted as a wave traveling in one direction since it consists of two terms, one traveling in one direction, and the other in the opposite direction. Nevertheless, a and b can be considered as waves associated with the incident and reflected powers, respectively. Even when the fields travel in one direction, e.g., when $B = 0$, there is, in general, a power reflection $|b|^2$. Therefore, when Z_0 is complex and A is fixed, it is necessary to introduce some field reflection, i.e., a finite B , if the maximum power is to be transmitted.

On the other hand, in contrast to the definition in (2.102) where a and b represented power waves, another pair of waves $\alpha(z)$ and $b(z)$ can be defined by

$$\alpha(z) = \frac{1}{2} Z_0^{-1/2} (E_x + Z_0 H_y), \quad b(z) = \frac{1}{2} Z_0^{-1/2} (E_x - Z_0 H_y) \quad (2.106)$$

corresponding to (1.16). Then, we have

$$\alpha(z) = Z_0^{-1/2} A e^{-(\alpha + j\beta)z}, \quad b(z) = Z_0^{-1/2} B e^{(\alpha + j\beta)z} \quad (2.107)$$

which indicate that $\alpha(z)$ and $b(z)$ are the traveling waves in the positive and negative z directions, respectively. Now however, the power expression corresponding to (2.103) can no longer be obtained (see Problem 2.8). In other words, the traveling waves $\alpha(z)$ and $b(z)$ are not associated directly with the transmission power when Z_0 is complex.

As we mentioned before, there is another pair of vectors E_y and H_x in addition to the one we have discussed so far. For this pair, repeating the same procedure, we obtain

$$E_y = C e^{-(\alpha + j\beta)z} + D e^{(\alpha + j\beta)z} \quad (2.108)$$

$$H_x = -Z_0^{-1} \{ C e^{-(\alpha + j\beta)z} - D e^{(\alpha + j\beta)z} \} \quad (2.109)$$

corresponding to (2.95) and (2.96), respectively. Because of the negative sign in front of the right-hand side of (2.109), when $D = 0$, the Poynting vector $\mathbf{E} \times \mathbf{H}^*$ points in the positive z -direction indicating that power is being transmitted in the direction of propagation. The relation between E_y and $-H_x$ is exactly the same as the relation between E_x and H_y , hence these two

electromagnetic waves have identical properties except for their field directions.

A more general expression for a plane wave is obtained by the superposition of $\mathbf{i}_x E_x$ and $\mathbf{i}_y E_y$ for the electric field and that of $\mathbf{i}_x H_x$ and $\mathbf{i}_y H_y$ for the magnetic field. The transmission power due to this superposition is given by

$$\begin{aligned}\operatorname{Re}(\mathbf{E} \times \mathbf{H}^*) &= \operatorname{Re}\{(\mathbf{i}_x E_x + \mathbf{i}_y E_y) \times (\mathbf{i}_x H_x^* + \mathbf{i}_y H_y^*)\} \\ &= \operatorname{Re}(\mathbf{i}_z E_x H_y^*) + \operatorname{Re}(\mathbf{i}_z E_y H_x^*)\end{aligned}\quad (2.110)$$

The first and second terms on the right-hand side of (2.110) express the power transmitted by the pairs E_x, H_y , and E_y, H_x , respectively. Thus the total power is the sum of the powers of individual waves. In other words, each wave can be considered as carrying its own power independent of the other. One might think this is obvious, however, if $E_x = A_1 e^{-j\beta z}$, $H_y = Z_0^{-1} A_1 e^{-j\beta z}$ represent one electromagnetic wave and $E_x = A_2 e^{-j\beta z}$, $H_y = Z_0^{-1} A_2 e^{-j\beta z}$ another, then the superposition of these two waves gives $E_x = (A_1 + A_2) e^{-j\beta z}$, $H_y = Z_0^{-1} (A_1 + A_2) e^{-j\beta z}$, and the transmission power becomes

$$\begin{aligned}\operatorname{Re}(\mathbf{E} \times \mathbf{H}^*) &= \operatorname{Re}[\mathbf{i}_z (A_1 + A_2) \{Z_0^{-1} (A_1 + A_2)\}^*] \\ &= \operatorname{Re}\{\mathbf{i}_z (Z_0^{-1})^* |A_1|^2\} + \operatorname{Re}\{\mathbf{i}_z (Z_0^{-1})^* |A_2|^2\} \\ &\quad + \operatorname{Re}\{\mathbf{i}_z (Z_0^{-1})^* (A_1 A_2^* + A_2 A_1^*)\}\end{aligned}$$

The first term on the right-hand side is the power transmitted by the first wave alone and the second term that due to the second wave alone. The existence of the third term shows that the transmission power is not given by the superposition of the individual powers. Thus, the principle of superposition does not generally hold for power. The reason why the total power is given by the direct sum of the powers of two waves in (2.110) is that the electric fields and hence the magnetic fields of these two waves are orthogonal in space. Generalizing this concept of orthogonality, whenever there are two waves whose powers add when their fields are superposed, they are said to be orthogonal to each other. Examples will be given in Chapter 3.

Let us now turn our attention to the superposed electric field itself. From (2.95) and (2.108), we have

$$\mathbf{E} = (\mathbf{i}_x A + \mathbf{i}_y C) e^{-(\alpha + j\beta)z} + (\mathbf{i}_x B + \mathbf{i}_y D) e^{(\alpha + j\beta)z} \quad (2.111)$$

Since the two terms on the right-hand side represent waves propagating in the positive and negative z -directions, respectively, let us concentrate on the first term and investigate how it changes with time. To do so, we take the real

the corresponding wave is said to be linearly polarized. When $|A| = |C|$ and $\varphi_A = \varphi_c \pm 90^\circ$, the ellipse becomes a circle and we have a circularly polarized wave. It may also be clear from Fig. 2.23 that the rotational direction becomes either clockwise or counterclockwise depending on whether $\varphi_A < \varphi_c$ or $\varphi_c < \varphi_A$. If we move the observation point z in the direction of propagation, the size of the ellipse becomes smaller and the direction of the field vector at the same instant changes.

The wave propagating in the negative z -direction has identical properties, e.g., depending on the amplitude and phase relations between B and D , the locus of the tip of vector $\mathbf{E}$ becomes an ellipse, a straight line, or a circle. The superposition of these two waves propagating in the positive and negative z -directions produces a locus which can be an ellipse, a straight line, or a circle at a fixed distance z . However, if the observation point is moved, the locus changes not only its size but also its shape resulting in a rather complicated polarization which depends on z .

So far in this section, we have assumed $\mathbf{E}$ and $\mathbf{H}$ to be independent of x and y , and we assumed terms of the type $\partial/\partial x$ and $\partial/\partial y$ to be zero. Let us next study the same waves from a slightly different viewpoint using vector expression with σ , ϵ , and μ assumed constant and independent of the position in space, as before.

Taking $\nabla \times$ (2.79) and substituting (2.80) into this result, we eliminate $\mathbf{H}$:

$$\nabla \times \nabla \times \mathbf{E} = -j\omega\mu \nabla \times \mathbf{H} = -j\omega\mu(\sigma + j\omega\epsilon) \mathbf{E}$$

The left-hand side is equal to $\nabla(\nabla \cdot \mathbf{E}) - \nabla^2 \mathbf{E}$ and since $\nabla \cdot \mathbf{E}$ is equal to zero, as can be shown by taking the divergence of (2.80), the above equation can be rewritten in the form

$$\nabla^2 \mathbf{E} + (\omega^2 \epsilon \mu - j\omega\mu\sigma) \mathbf{E} = 0 \quad (2.113)$$

Comparing (2.113) with (2.91), we see that d^2/dz^2 and E_x in (2.91) have been replaced by ∇^2 and $\mathbf{E}$, respectively, however, they both belong to the same type of differential equation. Since the solution of (2.91) was obtained by assuming the functional form of $Ae^{-\gamma z}$, in the present case $\mathbf{E} = \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r}) = \mathbf{A} e^{-(h_x x + h_y y + h_z z)}$ should give a solution where $\mathbf{h}$ is a constant vector to be determined. Using the relation

$$\begin{aligned} \nabla^2 \mathbf{E} &= \mathbf{A} \{(\partial^2/\partial x^2) + (\partial^2/\partial y^2) + (\partial^2/\partial z^2)\} e^{-(h_x x + h_y y + h_z z)} \\ &= \mathbf{A} (h_x^2 + h_y^2 + h_z^2) e^{-(h_x x + h_y y + h_z z)} \\ &= \mathbf{A} (\mathbf{h} \cdot \mathbf{h}) \exp(-\mathbf{h} \cdot \mathbf{r}) \end{aligned}$$

(2.113) becomes

$$\mathbf{A} \{ \mathbf{h} \cdot \mathbf{h} + (\omega^2 \epsilon \mu - j\omega \mu \sigma) \} \exp(-\mathbf{h} \cdot \mathbf{r}) = 0$$

Therefore,

$$\mathbf{h} \cdot \mathbf{h} = -(\omega^2 \epsilon \mu - j\omega \mu \sigma) \quad (2.114)$$

is a necessary and sufficient condition for $\mathbf{E} = \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r})$ to be a solution of (2.113). It is also necessary to check whether or not this satisfies Maxwell's equation which is done by substituting this expression into (2.79) and (2.80). From (2.79), we have

$$\nabla \times \mathbf{E} = \nabla \times \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r}) = \nabla \exp(-\mathbf{h} \cdot \mathbf{r}) \times \mathbf{A} = -j\omega \mu \mathbf{H}$$

Since

$$\begin{aligned} \nabla \exp(-\mathbf{h} \cdot \mathbf{r}) &= \{ \mathbf{i}_x (\partial/\partial x) + \mathbf{i}_y (\partial/\partial y) + \mathbf{i}_z (\partial/\partial z) \} e^{-(h_x x + h_y y + h_z z)} \\ &= -(\mathbf{i}_x h_x + \mathbf{i}_y h_y + \mathbf{i}_z h_z) \exp(-\mathbf{h} \cdot \mathbf{r}) \\ &= -\mathbf{h} \exp(-\mathbf{h} \cdot \mathbf{r}) \end{aligned}$$

$\mathbf{H}$ is given by

$$\mathbf{H} = (j\omega \mu)^{-1} \mathbf{h} \times \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r}) = (j\omega \mu)^{-1} \mathbf{h} \times \mathbf{E} \quad (2.115)$$

Substituting this into (2.80), we obtain

$$\begin{aligned} \nabla \times \mathbf{H} &= \nabla \times (j\omega \mu)^{-1} \mathbf{h} \times \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r}) = \nabla \exp(-\mathbf{h} \cdot \mathbf{r}) \times (j\omega \mu)^{-1} (\mathbf{h} \times \mathbf{A}) \\ &= -(j\omega \mu)^{-1} \mathbf{h} \times \mathbf{h} \times \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r}) \\ &= (\sigma + j\omega \epsilon) \mathbf{E} = (\sigma + j\omega \epsilon) \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r}) \end{aligned}$$

This requires $\mathbf{A}$ to satisfy

$$\mathbf{h} \times \mathbf{h} \times \mathbf{A} = (\omega^2 \epsilon \mu - j\omega \mu \sigma) \mathbf{A}$$

The left-hand side is equal to $\mathbf{h}(\mathbf{h} \cdot \mathbf{A}) - (\mathbf{h} \cdot \mathbf{h}) \mathbf{A}$, and since $\mathbf{h} \cdot \mathbf{h}$ is given by (2.114), the above requirement reduces to

$$\mathbf{h} \cdot \mathbf{A} = 0 \quad (2.116)$$

It follows from the above discussion that $\mathbf{E} = \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r})$ satisfies Maxwell's equations provided that (2.114) and (2.116) are satisfied.

Both $\mathbf{A}$ and $\mathbf{h}$ can be complex vectors. For example, the directions of the real and imaginary parts of the vector $\mathbf{h}$ may not be the same, and, therefore, the direction of maximum attenuation may be different from the direction of maximum phase variation. For simplicity, however, let us confine ourselves to the case in which these two directions coincide and let $\mathbf{n}$ be a unit vector in that direction. Then,

$$\mathbf{h} = \mathbf{n}(\alpha + j\beta) \quad (2.117)$$

where α and β are real. Since $\exp(-\mathbf{h} \cdot \mathbf{r}) = \exp\{-(\alpha + j\beta) \mathbf{n} \cdot \mathbf{r}\}$, both $\mathbf{E}$ and $\mathbf{H}$ become constant over the surface where $\mathbf{n} \cdot \mathbf{r}$ is constant. As is illustrated in Fig. 2.24, a constant $\mathbf{n} \cdot \mathbf{r}$ surface is a plane and, hence, the electromagnetic wave under consideration is a plane wave. Substituting (2.117) into (2.114), we have

$$(\alpha + j\beta)^2 = -(\omega^2 \epsilon \mu - j\omega \mu \sigma)$$

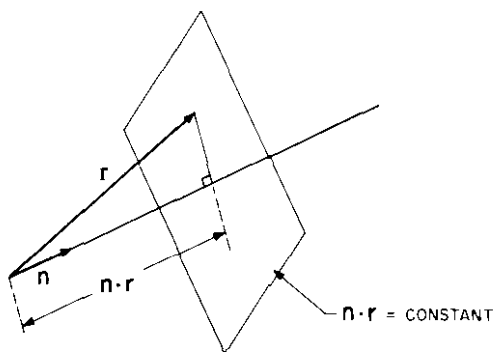


Fig. 2.24. Plane satisfying $\mathbf{n} \cdot \mathbf{r} = \text{constant}$.

This indicates that α and β are exactly the same as those obtained from (2.94). In addition, by comparing (2.116) and (2.117), it can be seen that

$$\mathbf{n} \cdot \mathbf{A} = 0 \quad (2.118)$$

since $(\alpha + j\beta)$ is constant and not equal to zero, and consequently, $\mathbf{A}$ is normal to the direction of propagation.

From (2.115), $\mathbf{H}$ is given by

$$\mathbf{H} = Z_0^{-1} \mathbf{n} \times \mathbf{E} \quad (2.119)$$

where Z_0 is the characteristic impedance defined by (2.97); and $\mathbf{H}$ is normal to both $\mathbf{n}$ and $\mathbf{E}$. We, therefore, see that the electric and magnetic fields are perpendicular to each other and at the same time both are normal to the direction of propagation. This corresponds to the result that E_x and H_y made one pair of vectors, and E_y and H_x another, while E_z and H_z were zero in the previous discussion.

As long as $\mathbf{A}$ and $\mathbf{h}$ satisfy (2.114) and (2.116), $\mathbf{E} = \mathbf{A} \exp(-\mathbf{h} \cdot \mathbf{r})$ gives a solution of Maxwell's equations. A superposition of many solutions of this type also constitutes a solution which will not be discussed here because of its complexity.

PROBLEMS

- 2.1 By decomposing vectors into their rectangular components, prove the following formulas:

$$\mathbf{A} \times (\mathbf{B} \times \mathbf{C}) = (\mathbf{A} \cdot \mathbf{C}) \mathbf{B} - (\mathbf{A} \cdot \mathbf{B}) \mathbf{C}$$

$$\nabla \cdot \nabla \times \mathbf{A} = 0$$

$$\nabla \times \nabla \phi = 0$$

- 2.2 Derive the expressions for $\nabla \cdot \mathbf{A}$, $\nabla \phi$, $\nabla \cdot \nabla \phi$ in the spherical coordinate system.

- 2.3 Prove in the cylindrical coordinate system, that $\nabla \times \mathbf{A}$ is given by

$$\nabla \times \mathbf{A} = \mathbf{i}_r \left(\frac{1}{r} \frac{\partial A_z}{\partial \theta} - \frac{\partial A_\theta}{\partial z} \right) + \mathbf{i}_\theta \left(\frac{\partial A_r}{\partial z} - \frac{\partial A_z}{\partial r} \right) + \mathbf{i}_z \left\{ \frac{1}{r} \frac{\partial}{\partial r} (r A_\theta) - \frac{1}{r} \frac{\partial A_r}{\partial \theta} \right\}$$

- 2.4 Prove that the integer n , satisfying $\nabla^2 r^n = 0$, is either 0 or -1 .

- 2.5 Prove that

$$\int_V \phi \nabla^2 \psi \, dv = - \int_V \nabla \phi \cdot \nabla \psi \, dv + \int_S \phi \nabla \psi \cdot \mathbf{n} \, dS$$

$$\int_S \phi \nabla^2 \psi \, dS = - \int_S \nabla \phi \cdot \nabla \psi \, dS + \int_L \phi \nabla \psi \cdot \mathbf{n} \, dl$$

- 2.6 Calculate the skin depth of copper at 4 GHz. The conductivity of copper is approximately $5.8 \times 10^7 \, \Omega/\text{m}$ and the permeability is 1. Also, calculate the characteristic impedance of copper at 4 GHz.

- 2.7 The tangential components of electric and magnetic fields must be continuous at an interface of two media. Assume that a plane wave is incident on an interface from the normal direction and calculate the power reflected relative to the incident power in terms of the dielectric constants and permeabilities of the media on both sides of the interface.

- 2.8 Prove that

$$|a(z)|^2 - |b(z)|^2 = |Z_0|^{-1} \operatorname{Re} \{ Z_0 E_x^* H_y \}$$

where $a(z)$ and $b(z)$ are defined by (2.106) and that the right-hand side is not equal to $\operatorname{Re} \{ E_x H_y^* \}$ when Z_0 is complex.

- 2.9 Show that Maxwell's equations (2.79) and (2.80) have no solution other than a trivial one when $\omega \neq 0$ and the electric and magnetic fields are functions of r only, i.e., distance from the origin. In other words, show that an electromagnetic wave with spherical symmetry cannot exist. (Hint: Use Gauss's and Stoke's theorems.)

CHAPTER 3

WAVEGUIDES

A cylindrical pipe designed to contain one or more propagating electromagnetic waves is called a waveguide. We shall discuss the theory of waveguides in detail in this chapter. Since waveguides require a treatment quite different from conventional circuit theory and radically new to some of us, we shall first introduce the particular case of rectangular waveguides. This will enable us to become acquainted with new terminology and some methods we shall use later. Following this will be an introduction to eigenvalue problems using an equation derived in the above discussion. Solutions of an eigenvalue problem are called the eigenfunctions, and these have wide applications in many branches of physics, particularly in acoustics and quantum mechanics. Suppose a function representing a linear phenomenon is to be determined, we can express it as a linear combination of eigenfunctions and determine the coefficients using appropriate equations governing the process. Once this expression is obtained, we interpret the phenomenon as the superposition of simple phenomena, each corresponding to an eigenfunction; the method is called the eigenfunction approach. For this approach to be useful, it is important to select eigenfunctions whose individual behavior is simple and well understood. Therefore, depending on the problem, a set of eigenfunctions is chosen to satisfy a certain eigenvalue problem closely related to the phenomenon under study. Whether or not the function to be determined can be expressed as a linear combination of eigenfunctions thus selected remains to be investigated. This, of course, depends on the eigenvalue problem used for their derivation, but when the linear combination is possible for any function of interest, the set of eigen-

functions is said to be complete. If the set does not have this property, the eigenfunction approach breaks down, even though all the coefficients could be determined, since the functions itself cannot be expressed as the linear combination in the first place. The completeness of eigenfunctions is discussed, therefore, as thoroughly as possible within the reach of elementary mathematics. This is followed by discussions on eigenfunctions for waveguides, their general theory, waveguide discontinuities, the effect of lossy walls, and waveguides with inhomogeneous media.

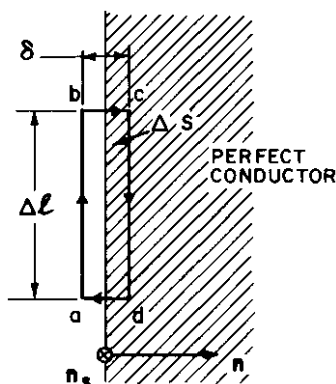
3.1 Perfect Conductors

Although actual conductors are not perfect, we shall, nevertheless, first study a straight waveguide bounded by a perfect conductor. To prepare for this discussion, let us consider what a perfect conductor is. A perfect conductor is an idealized medium within which the electric field is always equal to zero. A good conductor is generally considered as a medium whose conductivity σ is large. Therefore, it seems appropriate to define a perfect conductor as a medium whose conductivity is infinite. The conditions $\sigma \rightarrow \infty$ and $\mathbf{E} = 0$ can be considered equivalent to each other, using the following reasoning. The power consumed in a volume V inside a conductor is given by $\int_V \sigma \mathbf{E}^2 dv$. If σ increases indefinitely while $\mathbf{E}$ remains finite, the power consumption must approach infinity. However, since any realistic power source can deliver only a finite amount of power, it seems appropriate to assume that the condition $\sigma \rightarrow \infty$ corresponds to $\mathbf{E} \rightarrow 0$ in ordinary cases. Furthermore, as we discussed in Section 2.3, the amplitude of a plane wave diminishes quickly as it propagates in a conductor, and as σ increases, the attenuation becomes so large that an electromagnetic field inside a good conductor must become negligibly small. In the limit as $\sigma \rightarrow \infty$, the field must be zero.

Let us consider an interface between a perfect conductor and a medium with finite values of conductivity, dielectric constant and permeability. Taking a rectangular path as shown in Fig. 3.1, we integrate both sides of (2.79) over the area S of the rectangle. By Stokes's theorem, the left-hand side can be converted to a line integral around the closed path $a-b-c-d$:

$$\int_{abcd} \mathbf{E} \cdot d\mathbf{l} = -j\omega\mu \int_{\Delta S} \mathbf{H} \cdot \mathbf{n}_s dS$$

where $\mathbf{n}_s$ is the unit vector normal to ΔS . Letting the width δ of the rectangle diminish, only those contributions from $\overline{ab}$ and $\overline{cd}$ remain on the left-hand

Fig. 3.1. Integral contour $abcd$.

side, while the right-hand side becomes negligibly small since $\mathbf{H} \cdot \mathbf{n}_s$ stays finite and S approaches zero. Thus, we have

$$\int_{ab} \mathbf{E} \cdot d\mathbf{l} + \int_{cd} \mathbf{E} \cdot d\mathbf{l} = 0$$

Assuming the length Δl of the rectangle is small, this equation can be rewritten in the form

$$(\mathbf{E}_1 - \mathbf{E}_2) \cdot \mathbf{n} \times \mathbf{n}_s \Delta l = -\mathbf{n} \times (\mathbf{E}_1 - \mathbf{E}_2) \cdot \mathbf{n}_s \Delta l = 0$$

where $\mathbf{n}$ is the unit vector normal to the interface, as shown in Fig. 3.1, and $\mathbf{E}_1$ is the electric field just outside the conductor and $\mathbf{E}_2$ just inside the conductor. The orientation of the rectangle, and hence that of $\mathbf{n}_s$, is arbitrary as long as $\mathbf{n}_s$ lies in the interface. Therefore, from the above equation we obtain

$$\mathbf{n} \times (\mathbf{E}_1 - \mathbf{E}_2) = 0$$

However, since $\mathbf{E}_2 = 0$ by the definition of a perfect conductor, the electric field $\mathbf{E}_1$ just outside the conductor must satisfy

$$\mathbf{n} \times \mathbf{E}_1 = 0 \quad (3.1)$$

This indicates that the tangential component of an electric field at the surface of a perfect conductor is zero. Since $\mathbf{n} \times \mathbf{E} = 0$ at the surface, the the Poynting vector $\mathbf{E} \times \mathbf{H}^*$ has no component directed toward the inside of a perfect conductor, i.e., $\mathbf{n} \cdot \mathbf{E} \times \mathbf{H}^* = 0$. Therefore, all of the incident power is reflected back from the surface and none penetrates the perfect conductor.

Let us consider the magnetic field at the interface. It follows from (2.79) that $\mathbf{H}$ is zero inside a perfect conductor since $\mathbf{E}$ is equal to zero. If $\mathbf{H}$ is zero, (2.80) shows that not only $\mathbf{E}$ but also $\sigma\mathbf{E}$ is zero inside a perfect conductor. That is, no conduction current flows inside it even though σ is infinite. However, conduction current can flow at the interface. Let us integrate (2.80) over the rectangular area shown in Fig. 3.1. By Stokes's theorem we obtain

$$-\mathbf{n} \times (\mathbf{H}_1 - \mathbf{H}_2) \cdot \mathbf{n}_s \Delta l = \sigma \mathbf{E} \cdot \mathbf{n}_s \Delta S = \mathbf{i} \cdot \mathbf{n}_s \Delta S$$

where $\mathbf{H}_1$ is the magnetic field just outside the conductor and $\mathbf{H}_2$ just inside it. Since σ is infinite, $\mathbf{i} \Delta S = \sigma \mathbf{E} \cdot \Delta S$ is assumed to be finite in the limit of $\delta \rightarrow 0$. Let us define the surface current density $\mathbf{K}$ through

$$\mathbf{K} \Delta l = \lim_{\delta \rightarrow 0} \mathbf{i} \Delta S$$

Since $\mathbf{H}_2$ is equal to zero, the above equations show that

$$-\mathbf{n} \times \mathbf{H}_1 = \mathbf{K} \quad (3.2)$$

The tangential component of $\mathbf{E}$ vanishes at the interface; however, that of $\mathbf{H}$ usually does not, and the corresponding surface current flows perpendicular to it.

3.2 TE₁₀ Mode in Rectangular Waveguides

Let us first consider a straight waveguide with a rectangular cross section as shown in Fig. 3.2. We choose this particular model because of its practical value and its simplicity in mathematical treatment. The inside medium is assumed to be homogeneous and its conductivity equal to zero. Furthermore, it is assumed that the electric field has y -component (E_y) only, i.e., $E_x = E_z = 0$. Under these conditions, Maxwell's equations become

$$(\partial H_z / \partial y) - (\partial H_y / \partial z) = 0 \quad (3.3)$$

$$(\partial H_x / \partial z) - (\partial H_z / \partial x) = j\omega\epsilon E_y \quad (3.4)$$

$$(\partial H_y / \partial x) - (\partial H_x / \partial y) = 0 \quad (3.5)$$

$$-(\partial E_y / \partial z) = -j\omega\mu H_x \quad (3.6)$$

$$0 = -j\omega\mu H_y \quad (3.7)$$

$$(\partial E_y / \partial x) = -j\omega\mu H_z \quad (3.8)$$

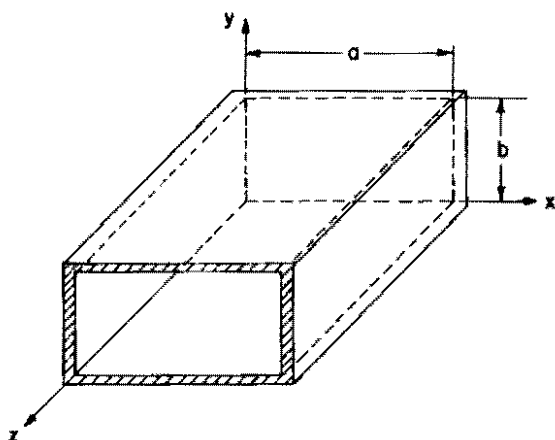


Fig. 3.2. Rectangular waveguide.

From (3.7), we obtain $H_y = 0$. Substitution of this result into (3.3) and (3.5) gives $\partial H_z / \partial y = \partial H_x / \partial y = 0$. Since $\partial H_y / \partial y = 0$, the magnetic field $\mathbf{H}$ does not change in the y -direction. Next, differentiating (3.6) and (3.8) with respect to z and x , respectively, and substituting the results into (3.4), we obtain the wave equation for E_y ,

$$\frac{\partial^2 E_y}{\partial x^2} + \frac{\partial^2 E_y}{\partial z^2} + \omega^2 \epsilon \mu E_y = 0 \quad (3.9)$$

Thus, the problem is reduced to the solution of this equation under appropriate boundary conditions given by $E_y = 0$ at $x = 0$ and $x = a$ which ensure that the tangential field component at the surface of the perfect conductor is equal to zero. Since $\mathbf{E} = \mathbf{i}_y E_y$ has no tangential components at $y = 0$ and $y = b$, the boundary condition there is automatically satisfied. Once E_y is obtained, H_x and H_z can be calculated from (3.6) and (3.8), respectively. Therefore, all the nonzero components of the electromagnetic field can be determined.

Let us consider particular solutions of (3.9) which vary with z as $e^{-\gamma z}$. Differentiation with respect to z therefore becomes equivalent to the multiplication by $-\gamma$. Thus, the wave equation and the boundary conditions become

$$\frac{\partial^2 E_y}{\partial x^2} + (\gamma^2 + \omega^2 \epsilon \mu) E_y = 0 \quad (3.10)$$

$$E_y = 0 \quad (x = 0, x = a) \quad (3.11)$$

Since (3.10) can be considered as an ordinary linear differential equation, the solutions must be exponential, sine or cosine functions of x . In order to satisfy the boundary conditions at $x=0$ and $x=a$, $\sin(n\pi x/a)$ must be used where n is an integer. Since (3.10) does not specify variation in the y -direction, we must take

$$E_y = A_n \sin(n\pi x/a) f(y) e^{-\gamma z}$$

as a trial function for the solution where $f(y)$ is an arbitrary function of y only. Substituting this expression into (3.6) and (3.8) and remembering that $\mathbf{H}$ has no variation in the y -direction, $f(y)$ is found to be a constant. Therefore, we have

$$E_y = A_n \sin(n\pi x/a) e^{-\gamma z} \quad (3.12)$$

Substituting this into (3.10), we obtain the necessary and sufficient condition for (3.12) to be a solution of (3.10), i.e.,

$$-(n\pi/a)^2 + \gamma^2 + \omega^2 \epsilon \mu = 0$$

or, equivalently,

$$\gamma = \pm \{(n\pi/a)^2 - \omega^2 \epsilon \mu\}^{1/2} \quad (3.13)$$

Substituting (3.12) into the original simultaneous equations (3.3) to (3.8), we find no contradiction regardless of the value of integer n in (3.13). This shows that E_y in (3.12), indeed, gives a solution for Maxwell's equations provided that γ satisfies (3.13).

Since the simultaneous equations (3.3) to (3.8) together with the boundary conditions (3.11) are linear, any linear combination of the functions on the right-hand side of (3.12) with various values of n should also represent a possible E_y . However, let us concentrate on individual terms. For each n , (3.12) gives a possible configuration of the electric field and hence that of the electromagnetic field in the waveguide. Such a configuration of the electromagnetic field is called a mode. For each mode, as ω is increased, γ always becomes purely imaginary beyond a particular value $\omega_c = (\epsilon \mu)^{-1/2} (n\pi/a)$. This indicates that the field as a whole begins to propagate through the waveguide above the frequency $f_c = \omega_c/2\pi$, where f_c is called the cutoff frequency of the mode. The free space cutoff wavelength, λ_c , corresponding to f_c is given by $2a/n$. Below f_c , γ is real, showing that the amplitude of the electromagnetic field changes exponentially with z . If a mode is excited below the cutoff frequency in a certain section of the waveguide, then the amplitude cannot increase with distance away from that section. Therefore, the sign of γ must be selected so that the wave amplitude decreases with distance in

both directions. It becomes positive on the positive side and negative on the negative side of the source of excitation.

As the frequency is raised, the first mode to start propagating corresponds to $n = 1$. In practice, if several modes propagate simultaneously, the behavior of the waveguide becomes complicated and hard to control. The cross-sectional dimensions are generally chosen so that the mode corresponding to $n = 1$ becomes the only propagating mode with the remaining modes in the cutoff region. Since no z -component (i.e., longitudinal component) exists in the electric field of the basic mode, it is called a transverse electric mode or, in short, a TE mode. For $n = 1$, since the y -directed field intensity has one maximum in the x -direction and none in y , the subscript 1,0 is attached, and the mode is identified as the TE₁₀ mode in the rectangular waveguide. This particular mode is sometimes called a dominant mode since it can exist when all other modes are in the cutoff condition, and they become negligibly small beyond a certain distance from their source of excitation.

The phase constant of the TE₁₀ mode is given by

$$\beta = \{\omega^2 \epsilon \mu - (\pi/a)^2\}^{1/2} = 2\pi\lambda^{-1} \{1 - (\lambda/\lambda_c)^2\}^{1/2} \quad (3.14)$$

where $\gamma = \pm j\beta$ and (3.13) are used. The corresponding wavelength in the waveguide becomes

$$\lambda_g = 2\pi/\beta = \lambda \{1 - (\lambda/\lambda_c)^2\}^{-1/2} \quad (3.15)$$

When the free space wavelength approaches λ_c , λ_g becomes longer and at $\lambda = \lambda_c$, it becomes infinite. Beyond λ_c , the field ceases to propagate.

All the nonzero components of the electromagnetic field are obtained from (3.12), (3.6) and (3.8). Thus, for the TE₁₀ mode, we have

$$\begin{aligned} E_x &= 0 \\ E_y &= (Ae^{-j\beta z} + Be^{j\beta z}) \sin(\pi x/a) \\ E_z &= 0 \\ H_x &= -(Z_0^{-1}) (Ae^{-j\beta z} - Be^{j\beta z}) \sin(\pi x/a) \\ H_y &= 0 \\ H_z &= j(\pi/a) (\omega\mu)^{-1} (Ae^{-j\beta z} + Be^{j\beta z}) \cos(\pi x/a) \end{aligned} \quad (3.16)$$

where A and B are constants and

$$Z_0 = (\mu/\epsilon)^{1/2} \{1 - (\lambda/\lambda_c)^2\}^{-1/2} \quad (3.17)$$

Let us investigate the field configuration concentrating on the terms with A only, since the terms with B give identical field configurations except that

the direction of the propagation is opposite. The field repeats itself with a period of λ_g in the z -direction. In the x -direction, E_y shows a sinusoidal variation with its maximum at the center ($x = a/2$) of the waveguide while H_x has the same variation except for the opposite sign. On the other hand, H_z is a cosine function of x with its zero at the center and with a phase differing 90° from E_y or H_x . Therefore, the electric and magnetic field configurations must look like those in Figs. 3.3(a) and (b), respectively.

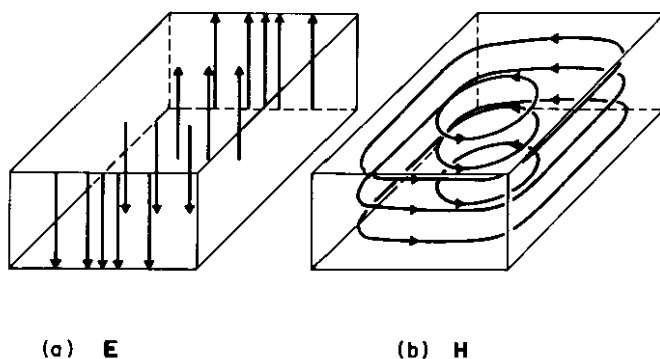


Fig. 3.3. Electric and magnetic fields in waveguide.

The vector $\mathbf{H}$ encircles the point where $\mathbf{E} = 0$ or, equivalently, where the displacement current $\epsilon \partial \mathbf{E} / \partial t$ becomes maximum. Figure 3.3 shows the patterns over a half wavelength in the z -direction at a fixed instant of time. The other half wavelength is identical except for the direction of the arrowheads. Those patterns move in the z -direction with time travelling with a phase velocity $v_p = \omega / \beta$.

Let us next consider the surface current on the waveguide wall, from (3.2). This is perpendicular to the magnetic field and hence should look like Fig. 3.4. If we remember that this pattern moves in the z -direction with time, at a fixed point on the wall, the surface current changes its magnitude and direction with time. At the center of the broad surface where $x = a/2$, the transverse current always remains zero. Therefore, the waveguide can either be split in two or a longitudinal slot can be cut along the center line without disturbing the inside fields appreciably.

To excite the TE_{10} mode, the center conductor of a coaxial line can be extended into the waveguide in the form of an antenna from the broad surface while its outer conductor is short-circuited to the wall. A short-

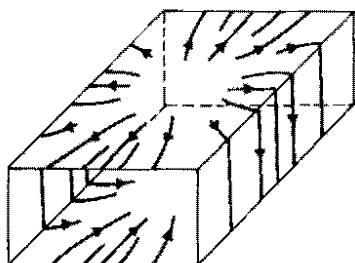


Fig. 3.4. Surface current on wall.

circuiting plunger is usually inserted in the waveguide behind the antenna to prevent power going in the other direction. This type of antenna excites a number of other modes, but if the TE_{10} mode is the only propagating mode, this can be launched efficiently by adjusting the size of the antenna as well as the position of the plunger.

A comparison of E_y , $-H_x$ in (3.16) with the voltage (1.9) and the current (1.14) on a transmission line shows that they are identical except for the factor $\sin(\pi x/a)$. As far as longitudinal variations are concerned, the electric and magnetic fields vary exactly in the same way as V and I on a transmission line with characteristic impedance Z_0 and phase constant β . Thus, introducing V and I , we can write the fields as follows:

$$E_y = KV \sin(\pi x/a) \quad (3.18)$$

$$-H_x = KI \sin(\pi x/a) \quad (3.19)$$

where K is a constant. Since the behavior of V and I can be studied on the Smith chart, the variation of E_y and $-H_x$ in the z -direction can also be investigated on the same chart. In the theory of transmission lines, the variation of the impedance Z defined by the ratio of V to I is obtained from the Smith chart. The corresponding quantity in the waveguide is the impedance defined by

$$Z = -(E_y/H_x) \quad (3.20)$$

The power transmitted in the positive z -direction is given by the real part of the integration of Poynting vector over the waveguide cross section,

$$\begin{aligned} \int \mathbf{E} \times \mathbf{H}^* \cdot \mathbf{k} \, dS &= \int E_y (-H_x^*) \, dS \\ &= \int \{\sin(\pi x/a)\}^2 K K^* V I^* \, dS = \frac{1}{2} ab K K^* V I^* \end{aligned}$$

where $\mathbf{k}$ is a unit vector in the z -direction. If K is defined as

$$K = \sqrt{2}(ab)^{-1/2} \quad (3.21)$$

then the power is given simply by $\text{Re}\{VI^*\}$ which is equal to the transmission power on the transmission line. It is worth mentioning that (3.21) is not the only condition for the power to be expressed as $\text{Re}\{VI^*\}$. Since the power relation specifies the magnitude of K only, the phase of K is still arbitrary. This phase is set equal to zero to avoid unnecessary complications. If $V' = nV$ and $I' = I/n$, where n is a positive constant, are substituted into (3.18) and (3.19), a comparison with (3.16) shows that V' and I' represent the voltage and current on a transmission line with the characteristic impedance $n^2 Z_0$ and the phase constant β . Furthermore, the transmission power in the waveguide is given by $\text{Re}\{V'I^*\}$ while the ratio of E_y to $-H_x$ is given by Z'/n^2 , where Z' is the impedance defined by the ratio of V' and I' . Therefore, V' and I' can be adopted equally well to represent the variation of the electric and magnetic fields inside the waveguide. It follows from this argument that the value of the characteristic impedance of the transmission line is not uniquely determined for representing the waveguide. In particular, the characteristic impedance $n^2 Z_0$ can be chosen to be unity in which case the difference between the normalized impedance and actual impedance disappears.

Although there is a one-to-one correspondence between the longitudinal variations of the transverse electric and magnetic fields of one mode in the waveguide and the variations of the voltage and current on the transmission line, the phase constant and the characteristic impedance considered above change with ω quite differently from those of the transmission line discussed in Section 1.1. There, Z_0 was independent of ω while β was proportional to ω . Since β given by (3.14) is not proportional to ω , the phase velocity $v_p = \omega/\beta$ changes with ω . Two waves, one at frequency ω and the other at $\omega + \Delta\omega$, have different velocities. When they are superposed, the velocity of the envelope is, therefore, different from either of their individual phase velocities. For simplicity, let us consider a superposition of two waves with the same amplitude. It is given by

$$\begin{aligned} & A \cos(\beta z - \omega t) + A \cos\{(\beta + \Delta\beta) z - (\omega + \Delta\omega) t\} \\ &= 2A \cos\{\tfrac{1}{2}(\Delta\beta z - \Delta\omega t)\} \cos\{(\beta + \tfrac{1}{2}\Delta\beta) z - (\omega + \tfrac{1}{2}\Delta\omega) t\} \end{aligned} \quad (3.22)$$

The first cosine term $\cos\{\tfrac{1}{2}(\Delta\beta z - \Delta\omega t)\}$ represents the envelope. Repeating a discussion similar to the one used for the phase velocity in Section 1.1,

the velocity of the envelope is found to be $\Delta\omega/\Delta\beta$. The last cosine term in (3.22) shows that the wave inside the envelope travels with the mean phase velocity of the two individual waves. If we illustrate the superposed wave as shown in Fig. 3.5, the dotted line representing the envelope moves with a velocity different from that of the solid line representing the detail. If we watch the space between two adjacent nodes of the envelope, the detailed wave may appear to be generated at the left-hand side and to propagate toward the right. At the nodes indicated by A_1 and A_2 in Fig. 3.5,

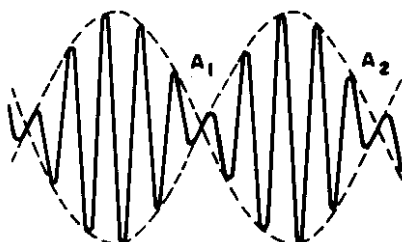


Fig. 3.5. Superposition of two waves with frequencies ω and $\omega + \Delta\omega$.

V and I , and, hence, the electric and magnetic fields are zero. It is, therefore, unlikely that electromagnetic energy is transferred in either direction across these points. The energy stored between A_1 and A_2 must travel with the same velocity as the envelope and is that carried by the superposition of the two waves. In the limit when $\Delta\omega \rightarrow 0$, the two frequencies coincide with each other, and the velocity of the envelope becomes

$$v_g = \lim_{\Delta\omega \rightarrow 0} \frac{\Delta\omega}{\Delta\beta} = \frac{d\omega}{d\beta} \quad (3.23)$$

This is called the group velocity. Since the argument about the velocity of energy holds regardless of the value of $\Delta\omega$, it is considered to be valid in the limit when $\Delta\omega \rightarrow 0$. Hence, the group velocity is considered to be the velocity of energy carried by the electromagnetic wave at frequency ω .

Note that the above argument holds only when the system is lossless; if it is lossy, nodes like the ones shown in Fig. 3.5 may not exist. Even if they do exist, as $\Delta\omega$ approaches zero, the distance between nodes becomes long and the wave amplitude decreases with distance. In the limiting case in which $\Delta\omega \rightarrow 0$, it disappears before reaching the next node. Therefore, the above conclusion of energy velocity being equal to the group velocity can no

longer be reached. In other words, the group velocity defined by (3.23) generally does not represent the velocity of energy in lossy cases.

For the present waveguide, we have from (3.14)

$$v_p = (\omega/\beta) = f\lambda \{1 - (\lambda/\lambda_c)^2\}^{-1/2} = v_0 \{1 - (\lambda/\lambda_c)^2\}^{-1/2} \quad (3.24)$$

$$\begin{aligned} v_g = \frac{d\omega}{d\beta} &= \left(\frac{d\beta}{d\omega} \right)^{-1} = (\omega\epsilon\mu)^{-1} \{ \omega^2\epsilon\mu - (\pi/a)^2 \}^{1/2} \\ &= \frac{1}{v_p} \frac{1}{\epsilon\mu} = \frac{v_0^2}{v_p} = v_0 \{1 - (\lambda/\lambda_c)^2\}^{1/2} \end{aligned} \quad (3.25)$$

where

$$v_0 = (\epsilon\mu)^{-1/2}$$

As the free space wavelength λ approaches λ_c , the phase velocity v_p increases indefinitely. On the other hand, the group velocity v_g always stays below v_0 confirming a result of the theory of relativity: The velocity of substance or energy never exceeds the light velocity. From (3.24) and (3.25), we obtain an interesting relation between v_p and v_g ,

$$v_p v_g = v_0^2$$

This relation holds for lossless waveguides filled with homogeneous media; however, one should not jump to a quick conclusion that the product of phase and group velocities is always a constant. Let us assume the product is constant, then we obtain the relation

$$v_p v_g = (\omega/\beta) (d\omega/d\beta) = \text{const}$$

which gives the solution $\beta = \pm(c_1\omega^2 + c_2)^{1/2}$ where c_1 and c_2 are arbitrary constants. The β given by (3.14) corresponds to this expression with $c_1 = \epsilon\mu$ and $c_2 = -(\pi/a)^2$. This result is generally not true; for example, when the waveguide is lossy, and hence the product of v_p and v_g is not necessarily a constant.

Finally, let us consider briefly the effect of discontinuities in the waveguide. If there is a window in the waveguide as shown in Fig. 3.6, the tangential electric field must vanish on the conducting walls indicated by the shaded area. This boundary condition is not generally satisfied by the electromagnetic field given in (3.16). Therefore, a number of other modes are generated to satisfy it when all are added together. Since these additional modes are usually in the cutoff region, their amplitudes decrease with distance away from the window, and beyond a certain distance d , all the modes except the TE₁₀ mode practically disappear. Let A and B in Fig. 3.7

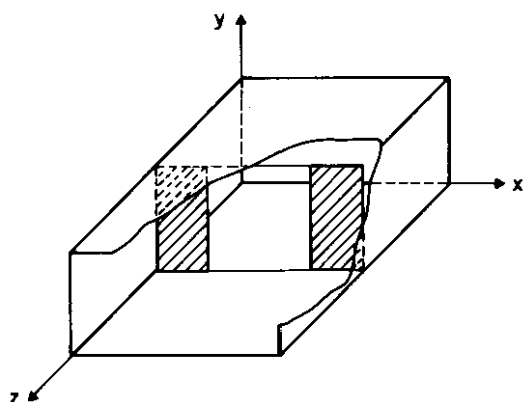


Fig. 3.6. Inductive window in waveguide.

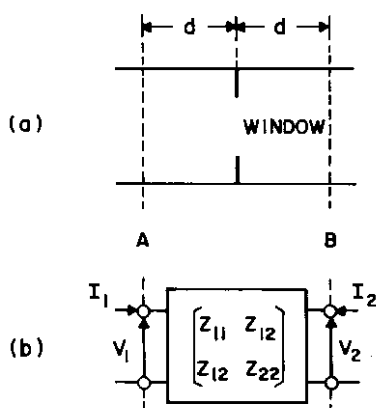


Fig. 3.7. Explanation of equivalent circuit of window (a) Waveguide top view; (b) Equivalent circuit.

be reference planes located as shown. Then as far as the region outside of *A* and *B* is concerned, the waveguide on each side can be represented by a transmission line corresponding to the TE_{10} mode. Since both Maxwell's equations and the boundary conditions are linear, there should be a linear relation between the E 's and H 's at *A* and *B* and, hence, between the V 's and I 's at the corresponding points on the transmission lines. This is expressed by

$$V_1 = Z_{11}I_1 + Z_{12}I_2, \quad V_2 = Z_{12}I_1 + Z_{22}I_2$$

where V_1 and V_2 are the voltages at A and B , respectively, and I_1 and I_2 are the corresponding currents. Thus, each discontinuity in the waveguide can be treated as a two-port network inserted in the corresponding transmission line. The coefficient of I_2 in the first equation and that of I_1 in the second equation are the same if the circuit is reciprocal. This point will be explained in Section 5.2 in detail.

When a thin window is placed perpendicular to the waveguide axis as shown in Fig. 3.6, it is known that the two-port network is equivalent to a length of transmission line with the same electrical length as the waveguide between A and B , shunted by a reactance at the position of the window, as shown in Fig. 3.8. Furthermore, if the window has the shape shown in the

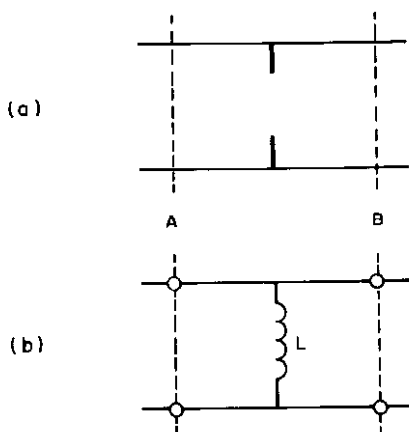


Fig. 3.8. Equivalent circuit of inductive window (a) Waveguide top view; (b) Corresponding equivalent circuit.

figure, the reactance is inductive. This may become plausible if we consider that the window introduces excess conduction current on the wall and, hence, increases the magnetic energy stored in its vicinity. A more detailed discussion will be given in Section 3.7. It is worth noting that the above equivalent circuit is valid only when the effect of the window on the region outside A and B is considered. If the equivalent circuit is used to discuss the region between A and B , wrong conclusions may be reached since the cutoff modes may not be adequately attenuated. For instance, suppose that an equivalent circuit of two inductive windows closely spaced to each other is expressed by two inductances inserted across the transmission line as shown in Fig. 3.9.

The value of each inductance is not generally equal to the value of the inductance for the same window situated an appreciable distance from other possible discontinuities. This is due to various interactions between the cutoff modes excited by the two windows, which do not exist in the case of a well-separated window.

The parameters of the two-port network Z_{11} , Z_{12} , and Z_{22} can be obtained by measuring the input impedance at A with various loads connected to B . The input impedance at A can be determined from a Smith chart, once the standing wave along the transmission line, and hence the relative magnitude of E_y in the waveguide has been measured as a function

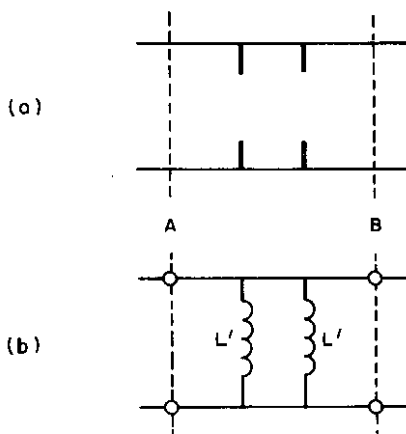


Fig. 3.9. Equivalent circuit of closely located inductive windows (a) Waveguide top view; (b) Corresponding equivalent circuit.

of z . This measurement is generally done using a movable probe inserted vertically into the waveguide through a longitudinal slot cut along the center line ($x = a/2$) of the broad surface. Since the slot does not interrupt the wall current as explained before, the TE_{10} mode disturbance is negligible. The probe couples to E_y only, and the output indicates its relative magnitude when detected and measured on a meter. From this, the standing wave ratio and the voltage minimum point can be determined, which are sufficient to obtain the impedance. A slotted waveguide with a movable probe for this purpose is called a standing wave detector.

The impedances required at B for the above measurement are provided by a movable short and a matched load. The former consists of a short-

circuiting metal plunger in the waveguide which completely reflects the incident wave. The position of the plunger and, hence, the phase of the reflected wave are adjustable, giving any impedance on the periphery of the Smith chart. The matched load usually consists of a resistive card longitudinally placed in the waveguide which completely absorbs the incident power, thus giving an impedance corresponding to the center of the Smith chart. The input end of the card is tapered to a point in order to eliminate reflections which might be caused by a sudden change in line properties. This technique of gradually changing line properties is often used to reduce reflections at discontinuities. That part of the line introduced for this purpose is called a tapered section.

3.3 Introductory Eigenvalue Problem

In the previous section, the problem of finding possible fields in the waveguide was reduced to that of solving the differential Eq. (3.10) under the boundary conditions (3.11). The constant $k^2 = \gamma^2 + \omega^2 \epsilon \mu$ had to be determined in order to make the solution satisfy the boundary conditions at $x = 0$ and $x = a$. As illustrated by that example, the problem of solving a differential equation that has a constant to be determined under appropriate boundary conditions is called an eigenvalue problem. The solutions are called the eigenfunctions and the corresponding constant for each eigenfunction is the eigenvalue. We know by substitution that the eigenfunctions and eigenvalues of (3.10) and (3.11) are given by

$$E_{yn} = A_n \sin(n\pi x/a), \quad k_n^2 = \gamma_n^2 + \omega^2 \epsilon \mu = (n\pi/a)^2$$

respectively. Since eigenvalue problems play an important role in the theory of waveguides, let us reconsider the same problem using a more general approach.

The k_n^2 's take discrete positive values in the above solutions, however, it is not clear whether or not other solutions exist in which the k_n^2 's are complex numbers. To investigate this point let E_{yn} be an eigenfunction of (3.10) and (3.11). By definition, we have

$$\frac{d^2 E_{yn}}{dx^2} + k_n^2 E_{yn} = 0 \quad (3.26)$$

$$E_{yn} = 0 \quad (x = 0, x = a) \quad (3.27)$$

Since E_{yn} may be a complex function, multiplying (3.26) by its conjugate, we

obtain

$$E_{yn}^* \frac{d^2 E_{yn}}{dx^2} + k_n^2 E_{yn}^* E_{yn} = \frac{d}{dx} \left(E_{yn}^* \frac{dE_{yn}}{dx} \right) - \frac{dE_{yn}^*}{dx} \frac{dE_{yn}}{dx} + k_n^2 E_{yn}^* E_{yn} = 0$$

Integrating this with respect to x from 0 to a , we have

$$\left[E_{yn}^* \frac{dE_{yn}}{dx} \right]_0^a - \int_0^a \frac{dE_{yn}}{dx}^2 dx + k_n^2 \int_0^a |E_{yn}|^2 dx = 0$$

where the limits of the integrals are from 0 to a , and $[]_0^a$ indicates the difference of values for the inside function at $x=0$ and $x=a$, e.g., $[f(x)]_0^a = f(a) - f(0)$. Since E_{yn} is equal to zero at $x=0$ and $x=a$, its conjugate is also zero at these two points, and the first term on the left hand side vanishes giving a simple expression for k_n^2 ;

$$k_n^2 = \frac{\int_0^a |dE_{yn}/dx|^2 dx}{\int_0^a |E_{yn}|^2 dx} \quad (3.28)$$

This expression shows that for a nonzero solution E_{yn} , k_n^2 can take neither negative nor complex values, which is the point we wanted to check. Now suppose that E_{yn} is a complex function, then since k_n^2 is real, the real and imaginary parts of E_{yn} satisfy (3.26), separately. Therefore, each E_{yn} can be considered as a real function without loss of generality.

To see why the eigenvalues are discrete, let us next investigate (3.26) and (3.27) with the restriction that E_y is real and k^2 positive. To satisfy the boundary condition at $x=0$, E_y starts from the origin as shown in Fig. 3.10. The solid and dotted lines illustrate the cases in which E_y starts into the positive or negative direction, respectively at $x=0$. Since E_y obeys the differential equation which can be rewritten in the form

$$\frac{1}{E_y} \frac{d^2 E_y}{dx^2} = -k^2$$

when E_y is positive, $d^2 E_y/dx^2$ must be negative, i.e., dE_y/dx must decrease with increasing x . On the other hand, when E_y is negative, $d^2 E_y/dx^2$ must be positive and dE_y/dx increases with increasing x . Therefore, E_y must be a convex function of x which tends to return to the base line $E_y=0$. If k^2 is small, this tendency to return to the base line is weak resulting in a nonzero value of E_y at $x=a$, as shown in Fig. 3.10(a). As k^2 increases, the curvature increases and a value of k^2 must exist which causes E_y to become exactly zero

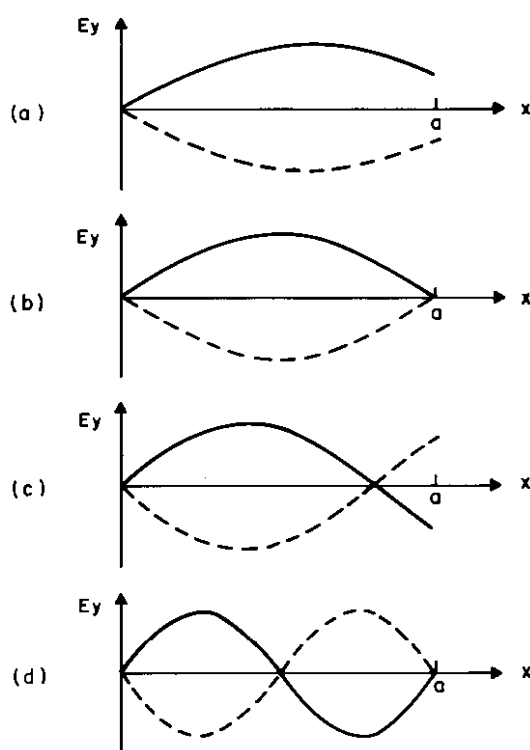


Fig. 3.10. Explanation of the existence of discrete eigenvalues.

at $x = a$, as shown in Fig. 3.10(b). This value of k^2 gives the first eigenvalue k_1^2 . The corresponding function E_y is the first eigenfunction E_{y1} . If k^2 increases further, E_y crosses the baseline and overswings at $x = a$, as shown in Fig. 3.10(c). Therefore, such a value of k^2 cannot be an eigenvalue since the boundary condition is not satisfied, but when k^2 increases further, another value occurs which causes E_y to return to the baseline exactly at $x = a$. This is shown in Fig. 3.10(d). The corresponding k^2 and E_y are the second eigenvalue k_2^2 and eigenfunction E_{y2} , respectively. In this way, k^2 can take only discrete values for E_y to satisfy the differential equation and the boundary conditions.

There is an important relation known as the orthogonality between any two different eigenfunctions having different eigenvalues and is given by

$$\int E_{yn} E_{ym} dx = 0 \quad (n \neq m) \quad (3.29)$$

where the integral is from $x = 0$ to $x = a$. The word orthogonality is employed here because of a similarity between (3.29) and the orthogonal relation between two vectors,

$$\mathbf{A} \cdot \mathbf{B} = A_x B_x + A_y B_y + A_z B_z = 0$$

A function defined over a finite interval can be represented conceptionally by a vector in an infinite dimensional abstract space whose components give the values of the function at each point in the interval. The scalar product between two such vectors can be expressed as the infinite summation of the products of corresponding components, which is proportional to the integral of the product of the corresponding functions. Therefore, the orthogonal relation is given by equating the integral to zero which is essentially done in (3.29). There is another reason for this terminology related to power which we shall discuss shortly.

The proof for (3.29) is as follows. Since E_{ym} is an eigenfunction, it satisfies

$$\frac{d^2 E_{ym}}{dx^2} + k_m^2 E_{ym} = 0 \quad (3.30)$$

Multiplying (3.26) and (3.30) by E_{yn} and E_{ym} , respectively, and integrating their difference from $x = 0$ to a , we have

$$\int \left(E_{ym} \frac{d^2 E_{yn}}{dx^2} - E_{yn} \frac{d^2 E_{ym}}{dx^2} \right) dx + (k_n^2 - k_m^2) \int E_{yn} E_{ym} dx = 0 \quad (3.31)$$

The first integral on the left-hand side can be transformed to

$$\begin{aligned} & \int \left\{ \frac{d}{dx} \left(E_{ym} \frac{dE_{yn}}{dx} \right) - \frac{dE_{ym}}{dx} \frac{dE_{yn}}{dx} \right\} - \left\{ \frac{d}{dx} \left(E_{yn} \frac{dE_{ym}}{dx} \right) - \frac{dE_{yn}}{dx} \frac{dE_{ym}}{dx} \right\} dx \\ &= \int \frac{d}{dx} \left(E_{ym} \frac{dE_{yn}}{dx} - E_{yn} \frac{dE_{ym}}{dx} \right) dx = \left[E_{ym} \frac{dE_{yn}}{dx} - E_{yn} \frac{dE_{ym}}{dx} \right]_0^a \end{aligned}$$

Since both E_{yn} and E_{ym} become zero at $x = 0$ and $x = a$, this is equal to zero and (3.31) reduces to

$$(k_n^2 - k_m^2) \int E_{yn} E_{ym} dx = 0 \quad (3.32)$$

$(k_n^2 - k_m^2)$ does not vanish by hypothesis, and hence (3.32) is equivalent to the orthogonality relation (3.29) which we wished to prove.

In our case, the eigenfunctions are sinusoidal and (3.29) corresponds to a

well-known formula

$$\int_0^a \sin(n\pi x/a) \sin(m\pi x/a) dx = 0 \quad (n \neq m) \quad (3.33)$$

This can be proved without relying on the above derivation as follows:

$$\begin{aligned} 2 \int \sin \frac{n\pi x}{a} \sin \frac{m\pi x}{a} dx &= \int \left\{ \cos \frac{(n-m)\pi x}{a} - \cos \frac{(n+m)\pi x}{a} \right\} dx \\ &= \left[\frac{a}{(n-m)\pi} \sin \frac{(n-m)\pi x}{a} \right]_0^a - \left[\frac{a}{(n+m)\pi} \sin \frac{(n+m)\pi x}{a} \right]_0^a = 0 \end{aligned}$$

In order to see the physical meaning of the orthogonality, let us calculate transmission power due to the superposition of two modes, n and m . From (3.6), the value of H_x corresponding to $E_y = A_n E_{yn}$ is given by

$$H_x = -\gamma_n (j\omega\mu)^{-1} A_n E_{yn}$$

Similarly, when $E_y = A_m E_{ym}$ then $H_x = -\gamma_m (j\omega\mu)^{-1} A_m E_{ym}$. The superposition of these two modes gives

$$E_y = (A_n E_{yn} + A_m E_{ym}), \quad H_x = -(j\omega\mu)^{-1} (A_n \gamma_n E_{yn} + A_m \gamma_m E_{ym})$$

The power in the z -direction is given by the integral of the Poynting vector over the cross section of the waveguide, i.e.,

$$\begin{aligned} \operatorname{Re} \int \mathbf{E} \times \mathbf{H}^* \cdot \mathbf{k} dS &= \operatorname{Re} \iint E_y (-H_x)^* dx dy \\ &= \operatorname{Re} \iint (A_n E_{yn} + A_m E_{ym}) (-j\omega\mu)^{-1} (A_n^* \gamma_n^* E_{yn} + A_m^* \gamma_m^* E_{ym}) dx dy \\ &= \operatorname{Re} \left\{ (-j\omega\mu)^{-1} \iint (|A_n|^2 \gamma_n^* E_{yn}^2 + A_n A_m^* \gamma_m^* E_{yn} E_{ym} \right. \\ &\quad \left. + A_m A_n^* \gamma_n^* E_{ym} E_{yn} + |A_m|^2 \gamma_m^* E_{ym}^2) dx dy \right\} \end{aligned}$$

Since the integral of $E_{yn} E_{ym}$ with respect to x vanishes from the orthogonality relation, the above expression reduces to

$$\operatorname{Re} \left\{ (-j\omega\mu)^{-1} \iint (|A_n|^2 \gamma_n^* E_{yn}^2 + |A_m|^2 \gamma_m^* E_{ym}^2) dx dy \right\} = P_n + P_m$$

where P_n and P_m are the transmission powers of the individual modes n and m , respectively. From this, we conclude that when the orthogonality relation (3.29) holds, the total transmission power is equal to the summation

of the powers transmitted separately by each mode. Referring to the discussion on the superposition of powers in Section 2.3, the above two modes n and m can be said to be orthogonal to each other, which is another reason for calling (3.29) the orthogonality relation.

Let us introduce an expression $k^2(E_y)$ given by

$$k^2(E_y) = \frac{\int (dE_y/dx)^2 dx - 2 [E_y (dE_y/dx)]_0^a}{\int E_y^2 dx} \quad (3.34)$$

where E_y is restricted to be real and the integrals apply to the interval from $x=0$ to $x=a$. For $k^2(E_y)$, the value is specified when a function of x , E_y , is given while the value of a function of x is specified when x is given. In general, a quantity whose value is fixed when a function is given is called a functional. Thus, $k^2(E_y)$ is a functional of E_y .

Let k^2 and $k^2 + \Delta k^2$ be values of $k^2(E_y)$ corresponding to E_y and $E_y + \Delta E_y$, respectively, where ΔE_y is a small change from E_y . Since $k^2 + \Delta k^2$ is given by the right-hand side of (3.34) when E_y is replaced everywhere by $E_y + \Delta E_y$, we have

$$(k^2 + \Delta k^2) \int (E_y + \Delta E_y)^2 dx = \int \{d(E_y + \Delta E_y)/dx\}^2 dx - 2 [(E_y + \Delta E_y) \{d(E_y + \Delta E_y)/dx\}]_0^a$$

where both sides have been multiplied by the denominator. Neglecting higher order terms of the small quantities Δk^2 and ΔE_y , the above equation reduces to

$$\begin{aligned} k^2 \int E_y^2 dx + 2k^2 \int \Delta E_y \cdot E_y dx + \Delta k^2 \int E_y^2 dx \\ = \int \left(\frac{dE_y}{dx} \right)^2 dx + 2 \int \left(\frac{d \Delta E_y}{dx} \cdot \frac{dE_y}{dx} \right) dx - 2 \left[E_y \frac{dE_y}{dx} \right]_0^a \\ - 2 \left[E_y \frac{d \Delta E_y}{dx} \right]_0^a - 2 \left[\Delta E_y \frac{dE_y}{dx} \right]_0^a \end{aligned} \quad (3.35)$$

It is worth mentioning that $d \Delta E_y/dx$ is also assumed to be small and its higher order terms are neglected in the above reduction. This assumption imposes a rather stringent requirement on possible ΔE_y 's; namely, when ΔE_y is said to be small in the following discussion, it means that not only the value of ΔE_y itself but also its derivative is small.

The first term on the left hand side of (3.35) cancels with the first and the third terms on the right-hand side because of (3.34). The second term on the right-hand side can be rewritten in the form

$$2 \int \left\{ \frac{d}{dx} \left(\Delta E_y \frac{dE_y}{dx} \right) - \Delta E_y \frac{d^2 E_y}{dx^2} \right\} dx = 2 \left[\Delta E_y \frac{dE_y}{dx} \right]_0^a - 2 \int \Delta E_y \frac{d^2 E_y}{dx^2} dx$$

Therefore, (3.35) becomes

$$2 \int \Delta E_y \left(\frac{d^2 E_y}{dx^2} + k^2 E_y \right) dx + 2 \left[E_y \frac{d \Delta E_y}{dx} \right]_0^a + \Delta k^2 \int E_y^2 dx = 0 \quad (3.36)$$

If E_y happens to be one of the eigenfunctions, say E_{yn} , the value of the expression $k^2(E_y)$ given by (3.34) is equal to the corresponding eigenvalue k_n^2 as we can see from (3.27) and (3.28). In this case, since E_{yn} satisfies both the differential equation and the boundary conditions, the first and second terms on the left-hand side of (3.36) vanish, and the first order variation Δk^2 of $k^2(E_y)$ due to a small variation ΔE_y from E_{yn} is thus equal to zero. That is, if a trial function E_y which is slightly different from a true eigenfunction is inserted in the expression $k^2(E_y)$, it gives an approximate value for the eigenvalue which is accurate up to the first order term of the difference ΔE_y . Conversely, if the first order variation Δk^2 is kept equal to zero for any small variation ΔE_y from a function E_y , then this function E_y is an eigenfunction for the following reason. Suppose (3.26) is not satisfied somewhere in the region under consideration, then ΔE_y can be chosen in such a way that it has the same sign as that of $d^2 E_y/dx^2 + k^2 E_y$ everywhere, provided the latter is not equal to zero. At the same time, $d \Delta E_y/dx$ can be chosen to be equal to zero at $x=0$ and $x=a$. Therefore, for this ΔE_y , the first term on the left-hand side of (3.36) becomes positive, the second term is zero and Δk^2 cannot be equal to zero. If, on the other hand, the differential equation is satisfied but the boundary condition at $x=0$ or $x=a$ is not satisfied, the sign of $d \Delta E_y/dx$ can be made the same as that of E_y at the boundary, giving nonzero Δk^2 again.

It follows from the above discussion that the following two statements are equivalent: (a) The first order variation Δk^2 of $k^2(E_y)$ is equal to zero for arbitrary ΔE_y 's whose values and derivatives are small; and (b) E_y is an eigenfunction. In other words, the eigenvalue problem is equivalent to that of finding E_y 's which give stationary values of $k^2(E_y)$ with respect to small variations in E_y . Such a functional $k^2(E_y)$ is called the variational expression for the eigenvalues.

Since $k^2(E_y) \geq 0$ for all functions satisfying the boundary conditions (3.27), there must be a function which minimizes $k^2(E_y)$. Let E_{y1} be this function, then, for any small variation ΔE_y satisfying the same boundary conditions, the first order variation Δk^2 from $k^2(E_{y1})$ is equal to zero. Since the second and third terms on the left-hand side of (3.36) are zero, E_{y1} must satisfy the differential equation. Otherwise, if the first term is made positive by choosing the sign of ΔE_y to be the same as that of $(d^2 E_{y1}/dx^2) + k^2 E_{y1}$ everywhere, then a contradiction is obtained. From this, it follows that E_{y1} thus chosen is an eigenfunction. Next, let E_{y2} be a function which satisfies the boundary conditions and at the same time minimizes $k^2(E_y)$ with the additional condition that it is orthogonal to E_{y1} . Then, for any small ΔE_y satisfying the boundary conditions, Δk^2 is again found to be zero as follows. Since an arbitrary function F can be written in the form

$$F = E_{y1} \frac{\int F \cdot E_{y1} dx}{\int E_{y1}^2 dx} + \left(F - E_{y1} \frac{\int F \cdot E_{y1} dx}{\int E_{y1}^2 dx} \right)$$

where the first term is proportional to E_{y1} and the second term is orthogonal to E_{y1} , ΔE_y can be decomposed into two parts, one proportional to E_{y1} and the other orthogonal. From the definition of E_{y2} , it is obvious that Δk^2 corresponding to that part of ΔE_y orthogonal to E_{y1} always vanishes. Let us, therefore, consider Δk^2 corresponding to the part of ΔE_y proportional to E_{y1} . We have

$$\begin{aligned} & \int E_{y1} \left(\frac{d^2 E_{y2}}{dx^2} + k_2^2 E_{y2} \right) dx \\ &= \int \left\{ \frac{d}{dx} \left(E_{y1} \frac{dE_{y2}}{dx} - E_{y2} \frac{dE_{y1}}{dx} \right) + E_{y2} \left(\frac{d^2 E_{y1}}{dx^2} + k_2^2 E_{y1} \right) \right\} dx \\ &= \left[E_{y1} \frac{dE_{y2}}{dx} - E_{y2} \frac{dE_{y1}}{dx} \right]_0^a + (k_2^2 - k_1^2) \int E_{y1} E_{y2} dx \end{aligned}$$

in which the first and second terms on the right-hand side disappear because of the boundary conditions for E_{y1} and E_{y2} and the orthogonality relation, respectively. Consequently, the first term on the left-hand side of (3.36) becomes zero for ΔE_y proportional to E_{y1} . The second term is zero from the boundary conditions for E_{y2} , thus leading to the conclusions that Δk^2 , corresponding to the part of ΔE_y proportional to E_{y1} , is equal to zero. Therefore, $\Delta k^2 = 0$

for any small ΔE_y satisfying the boundary conditions. Applying the same argument used for E_{y1} , E_{y2} is found to be another eigenfunction. Similarly, let E_{y3} be a function which satisfies the boundary conditions and minimizes $k^2(E_y)$ under two additional conditions of being orthogonal to E_{y1} and E_{y2} . An argument similar to the above shows that E_{y3} is also an eigenfunction. In this way, adding the orthogonality conditions one by one, one can find an infinite series of eigenfunctions $E_{y1}, E_{y2}, \dots$ with the corresponding eigenvalues satisfying $0 \leq k_1^2 \leq k_2^2 \leq \dots$, and k_n^2 thus obtained increases indefinitely with n , i.e.,

$$\lim_{n \rightarrow \infty} k_n^2 = \infty$$

A proof as to why k_n^2 increases with n which can be generalized to two and three dimensional cases will be given in Appendix I.

If we now assume the infinite growth of the eigenvalues, one of the most important properties of the set of eigenfunctions, called the completeness, can be derived as follows. Since an eigenfunction multiplied by a constant is still an eigenfunction with the same eigenvalue, let us first normalize every eigenfunction by multiplying with a proper constant, i.e., for every E_{yn} , the relation

$$\int E_{yn}^2 dx = 1 \quad (3.37)$$

is assumed to hold. In our case, $E_{yn} = \sqrt{2} a^{-1/2} \sin(n\pi x/a)$ is the appropriate form of normalized eigenfunctions. For convenience, let us call a function piecewise-continuous when it is continuous except at a finite number of points in the domain of interest, i.e., in the interval from $x=0$ to $x=a$ in the present case. We shall also call it square-integrable when the integral of its square exists over the same domain. Let a function f satisfying the boundary conditions have a square-integrable derivative, but otherwise, be arbitrary. We define f_N by

$$f_N = f - \sum_{n=1}^{N-1} A_n E_{yn} \quad (3.38)$$

where

$$A_n = \int f \cdot E_{yn} dx \quad (3.39)$$

Also, let a_N^2 be given by

$$a_N^2 = \int f_N^2 dx \quad (3.40)$$

From (3.38) and (3.39), using the orthogonality and normalization conditions, we have

$$\begin{aligned} a_N^2 &= \int (f - \sum A_n E_{yn})^2 dx \\ &= \int (f^2 - 2 \sum A_n f E_{yn} + \sum A_n^2 E_{yn}^2) dx = \int f^2 dx - \sum_{n=0}^{N-1} A_n^2 \end{aligned} \quad (3.41)$$

From the definitions of f_N and a_N , we obtain

$$\int (f_N/a_N)^2 dx = 1 \quad (3.42)$$

and

$$\int (f_N/a_N) E_{yn} dx = 0 \quad (n < N) \quad (3.43)$$

Since f_N/a_N is orthogonal to E_{yn} ($n < N$) and satisfies the boundary conditions, $k^2(f_N/a_N)$ cannot be less than k_N^2 . This is due to the fact that E_{yN} minimizes $k^2(E_y)$ to k_N^2 while at the same time it satisfies the boundary conditions and is orthogonal to each E_{yn} for which $n < N$. The denominator of $k^2(f_N/a_N)$ is unity from (3.42), and hence we have

$$\begin{aligned} k^2 \left(\frac{f_N}{a_N} \right) &= \frac{1}{a_N^2} \int \left(\frac{df_N}{dx} \right)^2 dx = \frac{1}{a_N^2} \int \left(\frac{df}{dx} - \sum A_n \frac{dE_{yn}}{dx} \right)^2 dx \\ &= \frac{1}{a_N^2} \int \left\{ \left(\frac{df}{dx} \right)^2 - 2 \sum A_n \frac{df}{dx} \frac{dE_{yn}}{dx} + \left(\sum A_n \frac{dE_{yn}}{dx} \right)^2 \right\} dx \\ &= \frac{1}{a_N^2} \int \left(\frac{df}{dx} \right)^2 dx - \frac{1}{a_N^2} 2 \sum A_n \left\{ \left[f \frac{dE_{yn}}{dx} \right]_0^a - \int f \frac{d^2 E_{yn}}{dx^2} dx \right\} \\ &\quad + \frac{1}{a_N^2} \sum \sum A_n A_m \left\{ \left[E_{yn} \frac{dE_{ym}}{dx} \right]_0^a - \int E_{yn} \frac{d^2 E_{ym}}{dx^2} dx \right\} \\ &= \frac{1}{a_N^2} \int \left(\frac{df}{dx} \right)^2 dx - \frac{1}{a_N^2} 2 \sum A_n^2 k_n^2 + \frac{1}{a_N^2} \sum A_n^2 k_n^2 \\ &= \frac{1}{a_N^2} \int \left(\frac{df}{dx} \right)^2 dx - \frac{1}{a_N^2} \sum_{n=1}^{N-1} A_n^2 k_n^2 \geq k_N^2 \end{aligned} \quad (3.44)$$

where the boundary conditions, (3.39) and the orthogonality conditions are used. Noting that the second term on the lefthand side of the inequality is positive, we have

$$a_N^{-2} \int (df/dx)^2 dx \geq k_N^2$$

or equivalently,

$$a_N^2 \leq k_N^{-2} \int (df/dx)^2 dx$$

It follows from this that

$$\lim_{N \rightarrow \infty} a_N^2 = 0$$

since $\int (df/dx)^2 dx$ has a finite value and k_N^2 increases indefinitely with N . From the definition of a_N , the above equation is equivalent to

$$\lim_{N \rightarrow \infty} \int \left(f - \sum_{n=1}^{N-1} A_n E_{y_n} \right)^2 dx = 0 \quad (3.45)$$

This means that an arbitrary function f which satisfies the boundary conditions and has a square-integrable derivative can be expanded in terms of the E_{y_n} 's which implies that the integral of the square of the difference between f and $\sum A_n E_{y_n}$ converges to zero.

Let F be an arbitrary function which is square integrable and piece-wise-continuous. Then F can be approximated by a function f which satisfies the boundary conditions and whose derivative is square-integrable, as shown in Fig. 3.11, so as to satisfy

$$\int (F - f)^2 dx < \varepsilon/4$$

where ε is a fixed, but arbitrarily small, positive value. On the other hand, from (3.45)

$$\int \left(f - \sum_{n=1}^{N-1} A_n E_{y_n} \right)^2 dx < \varepsilon/4$$

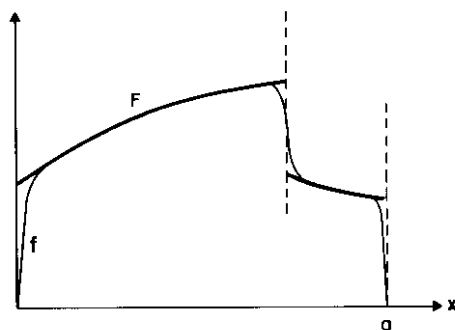


Fig. 3.11. Approximation of a discontinuous function by a continuous function with square-integrable derivative.

for a sufficiently large N . Therefore, we have

$$\begin{aligned} \int \left(F - \sum_{n=1}^{N-1} A_n E_{yn} \right)^2 dx &= \int (F - f + f - \sum A_n E_{yn})^2 dx \\ &\leq 2 \int (F - f)^2 dx + 2 \int (f - \sum A_n E_{yn})^2 dx \leq \varepsilon \end{aligned}$$

where use is made of the relation (2) in Appendix I. Since ε is arbitrary, the above relation shows that a piecewise continuous and square-integrable, but otherwise arbitrary, function F can be expanded in terms of the E_{yn} 's:

$$\lim_{N \rightarrow \infty} \int \left(F - \sum_{n=1}^{N-1} A_n E_{yn} \right)^2 dx = 0 \quad (3.46)$$

This property of the E_{yn} 's is called the completeness. The above relation is usually written in the form

$$F = \sum_{n=1}^{\infty} A_n E_{yn} \quad (3.47)$$

where

$$A_n = \int_0^a F \cdot E_{yn} dx \quad (3.48)$$

The reason that (3.39) can be rewritten in the form of (3.48) is easily seen from the following:

$$\begin{aligned} \left\{ \int F \cdot E_{yn} dx - \int f \cdot E_{yn} dx \right\}^2 &= \left\{ \int (F - f) \cdot E_{yn} dx \right\}^2 \\ &\leq \int (F - f)^2 dx \int E_{yn}^2 dx \leq \varepsilon/4 \end{aligned}$$

where use is made of (1) in Appendix I and the normalization condition for E_{yn} .

In the present case of the eigenvalue problem under consideration (3.26) and (3.27), the E_{yn} 's are sinusoidal functions and (3.47) is the Fourier expansion of F :

$$F = \sum_{n=1}^{\infty} A_n \sqrt{2} a^{-1/2} \sin(n\pi x/a)$$

where

$$A_n = \int_0^a F \cdot \sqrt{2} a^{-1/2} \sin(n\pi x/a) dx$$

Although we wrote (3.47) in the form of an ordinary equation, the real

meaning of the equality is given by (3.46) in which the equality does not mean that the value of the left-hand side of (3.47) is equal to the value of the right at each point in the domain of interest. For instance, it is well known that the Fourier expansion gives the mean value at each discontinuity point. This restriction does not bother us, however, since any measured value of a physical quantity at a point is actually the average over a certain region in the vicinity of that point. It is immaterial what value is given by the functions at each point as long as the average value over a neighborhood stays the same. In this sense, both sides of (3.47) are equal.

With this understanding of the meaning of the equality, let us next study how to use it. Let the expansion of F and G be given by $\sum A_n E_{yn}$ and $\sum B_n E_{yn}$, respectively. If F and G are equal in the sense that $\int (F - G)^2 dx = 0$, then $A_n = B_n$ for each n . The proof is as follows. For a given ε no matter how small it may be, the following inequality can be satisfied with a sufficient large N where Lemma 2 in Appendix I is used twice for the second inequality.

$$\begin{aligned} & \int \left(\sum_{n=1}^{N-1} A_n E_{yn} - \sum_{n=1}^{N-1} B_n E_{yn} \right)^2 dx \\ & \leq \int (F - G + \sum A_n E_{yn} - F + G - \sum B_n E_{yn})^2 dx \\ & \leq 2 \int (F - G)^2 dx + 4 \int (F - \sum A_n E_{yn})^2 dx + 4 \int (G - \sum B_n E_{yn})^2 dx \\ & \leq \varepsilon \end{aligned}$$

However, the left-hand side is equal to $\sum_{n=1}^{N-1} (A_n - B_n)^2$ because of the orthogonality between the E_{yn} 's. It follows from this that $A_n = B_n$ for each n , otherwise ε cannot be arbitrarily small. Conversely, if $A_n = B_n$ for every n ,

$$\begin{aligned} \int (F - G)^2 dx &= \int (F - \sum A_n E_{yn} + \sum B_n E_{yn} - G)^2 dx \\ &\leq 2 \int (F - \sum A_n E_{yn})^2 dx + 2 \int (G - \sum B_n E_{yn})^2 dx \\ &\leq \varepsilon \end{aligned}$$

Therefore $F = G$ in the sense that $\int (F - G)^2 dx = 0$. The same conclusion can be obtained by a slightly less rigorous but more popular method. Disregarding the real meaning of the equality between series expansions, first equate the two expansion forms as $\sum A_n E_{yn} = \sum B_n E_{yn}$, next multiply both sides by E_{yn} and integrate with respect to x over the domain of interest from $x = 0$ to $x = a$. We are then left with only one term, $A_n = B_n$, because

of the orthogonality of the E_{yn} 's. Conversely, if $A_n = B_n$ for every n , obviously $\sum_{n=1}^N A_n E_{yn} = \sum_{n=1}^N B_n E_{yn}$ for all N , which means $F = G$.

Another interesting formula is given by

$$\int F \cdot G \, dx = \sum_{n=1}^{\infty} A_n B_n \quad (3.49)$$

To prove this, let us calculate the square of

$$\int \left(F - \sum_{n=1}^{N-1} A_n E_{yn} \right) \left(G - \sum_{n=1}^{N-1} B_n E_{yn} \right) dx$$

The result is

$$\begin{aligned} & \left\{ \int (F - \sum A_n E_{yn}) (G - \sum B_n E_{yn}) \, dx \right\}^2 \\ &= \left\{ \int F \cdot G \, dx - \sum A_n \int G \cdot E_{yn} \, dx - \sum B_n \int F \cdot E_{yn} \, dx + \sum A_n B_n \right\}^2 \\ &= \left\{ \int F \cdot G \, dx - \sum A_n B_n \right\}^2 \end{aligned}$$

However, because of Lemma 1 in Appendix I, the left-hand side is smaller than

$$\int \left(F - \sum_{n=1}^{N-1} A_n E_{yn} \right)^2 dx \cdot \int \left(G - \sum_{n=1}^{N-1} B_n E_{yn} \right)^2 dx$$

which can be made arbitrarily small. Thus, for a given $\epsilon > 0$, no matter how small it may be,

$$\left\{ \int F \cdot G \, dx - \sum_{n=1}^{N-1} A_n B_n \right\}^2 < \epsilon$$

This inequality can be satisfied by making N sufficiently large, which is equivalent to the desired formula (3.49). In particular, when $F = G$, we have

$$\int F^2 \, dx = \sum_{n=1}^{\infty} A_n^2$$

The same conclusion can be obtained by the following less rigorous method:

$$\int F \cdot G \, dx = \int \sum_{n=1}^{\infty} A_n E_{yn} \cdot \sum_{n=1}^{\infty} B_n E_{yn} \, dx = \sum_{n=1}^{\infty} A_n B_n$$

where the orthogonality and normalization conditions between the E_{yn} 's are used.

Finally, it is worth mentioning that, when F is complex, the real and imag-

inary parts can be expanded separately and then added together to obtain exactly the same formula as (3.47) and (3.48). In this case, (3.46) is replaced by

$$\lim_{n \rightarrow \infty} \int \left| F - \sum_{n=1}^{N-1} A_n E_{yn} \right|^2 dx = 0$$

Similarly, (3.49) holds even if F and G are complex.

3.4 Eigenfunctions for Waveguides

In Section 3.2, we investigated properties of electromagnetic waves in rectangular waveguides, assuming that the fields varied exponentially in the longitudinal direction and that E_y was the only nonzero component of the electric field. In this section, eliminating the second assumption, we shall set up an eigenvalue problem for the transverse electric field in a straight waveguide with arbitrary cross section. The solutions of this problem play an important role in the theory of waveguides to be discussed in the next section.

Because of the particular geometry of a waveguide extended in one direction, for example in the z -direction, and limited in other directions, it is convenient to first separate the fields into their longitudinal and transverse components:

$$\mathbf{E} = (\mathbf{E}_t + \mathbf{k}E_z) e^{-\gamma z} \quad (3.50)$$

$$\mathbf{H} = (\mathbf{H}_t + \mathbf{k}H_z) e^{-\gamma z} \quad (3.51)$$

where the exponential variation with respect to z is written explicitly so that $\mathbf{E}_t$, $\mathbf{H}_t$, E_z , and H_z become independent of z . Vector $\mathbf{k}$ is a unit vector in the z direction while $\mathbf{E}_t$, $\mathbf{H}_t$ are the transverse components of the electric and magnetic fields, respectively.

Let us substitute (3.50) and (3.51) into Maxwell's equations

$$\nabla \times \mathbf{H} = j\omega\epsilon\mathbf{E} \quad (3.52)$$

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H} \quad (3.53)$$

where σ is assumed to be zero. First, substituting into (3.52), we have

$$\nabla \times (\mathbf{H}_t + \mathbf{k}H_z) e^{-\gamma z} = j\omega\epsilon(\mathbf{E}_t + \mathbf{k}E_z) e^{-\gamma z} \quad (3.54)$$

The left-hand side of (3.54) consists of two terms,

$$\begin{aligned} \nabla \times \mathbf{H}_t e^{-\gamma z} &= e^{-\gamma z} \nabla \times \mathbf{H}_t + \nabla e^{-\gamma z} \times \mathbf{H}_t \\ &= e^{-\gamma z} \nabla \times \mathbf{H}_t - \mathbf{k}\gamma e^{-\gamma z} \times \mathbf{H}_t \end{aligned} \quad (3.55)$$

$$\begin{aligned}\nabla \times \mathbf{k} H_z e^{-\gamma z} &= e^{-\gamma z} \nabla H_z \times \mathbf{k} + \nabla e^{-\gamma z} \times \mathbf{k} H_z \\ &= -e^{-\gamma z} \mathbf{k} \times \nabla H_z - \gamma e^{-\gamma z} \times \mathbf{k} H_z\end{aligned}\quad (3.56)$$

where the ordinary rule for differentiating a product of functions is used. The second term on the right-hand side of (3.55) and the first term on the right-hand side of (3.56) are both perpendicular to $\mathbf{k}$. The second term on the right-hand side of (3.56) is equal to zero because of its form $\mathbf{k} \times \mathbf{k}$. In order to find the direction of $\nabla \times \mathbf{H}_t$, let us consider

$$\mathbf{k} \times (\nabla \times \mathbf{H}_t) = \nabla (\mathbf{k} \cdot \mathbf{H}_t) - (\mathbf{k} \cdot \nabla) \mathbf{H}_t$$

Since $\mathbf{H}_t$ has no $\mathbf{k}$ component, the first term on the right-hand side is equal to zero. Because $\mathbf{H}_t$ is independent of z and $\mathbf{k} \cdot \nabla = \partial/\partial z$, the second term also vanishes. Thus, $\mathbf{k} \times (\nabla \times \mathbf{H}_t)$ is found to be zero from which we can conclude that $\nabla \times \mathbf{H}_t$ is parallel to $\mathbf{k}$ and the first term on the right-hand side of (3.55) gives the $\mathbf{k}$ -component only. Equation (3.54) can now be decomposed into two equations, one for the $\mathbf{k}$ -component and the other for the transverse component:

$$\nabla \times \mathbf{H}_t = j\omega\epsilon\mathbf{k}E_z \quad (3.57)$$

$$\gamma\mathbf{k} \times \mathbf{H}_t + \mathbf{k} \times \nabla H_z = -j\omega\epsilon\mathbf{E}_t \quad (3.58)$$

Similarly, from (3.53) we have two equations

$$\nabla \times \mathbf{E}_t = -j\omega\mu\mathbf{k}H_z \quad (3.59)$$

$$\gamma\mathbf{k} \times \mathbf{E}_t + \mathbf{k} \times \nabla E_z = j\omega\mu\mathbf{H}_t \quad (3.60)$$

Next, applying $\nabla \cdot$ to (3.52), we obtain

$$\nabla \cdot \mathbf{E} = 0$$

since $\nabla \cdot \nabla \times$ is always equal to zero. A substitution of (3.50) into this equation gives

$$e^{-\gamma z} (\nabla \cdot \mathbf{E}_t + \nabla \cdot \mathbf{k} E_z) + (\mathbf{E}_t + \mathbf{k} E_z) \cdot \nabla e^{-\gamma z} = 0$$

Since E_z is independent of z and $\nabla \cdot \mathbf{k} = \partial/\partial z$, $\nabla \cdot \mathbf{k} E_z$ is therefore equal to zero. In addition $\nabla e^{-\gamma z}$ is equal to $\mathbf{k}(-\gamma)e^{-\gamma z}$, and since $\mathbf{E}_t$ is perpendicular to $\mathbf{k}$, $\mathbf{E}_t \cdot \nabla e^{-\gamma z}$ also vanishes. Thus, we are left with

$$\nabla \cdot \mathbf{E}_t = \gamma E_z \quad (3.61)$$

Similarly, we have from $\nabla \cdot$ (3.53)

$$\nabla \cdot \mathbf{H}_t = \gamma H_z \quad (3.62)$$

We shall now try to obtain an equation for $\mathbf{E}_t$ alone by eliminating $\mathbf{H}_t$, H_z , and E_z from the above equations. First, calculate $\nabla \times$ (3.59):

$$\begin{aligned}\nabla \times \nabla \times \mathbf{E}_t &= -j\omega\mu \nabla \times \mathbf{k}H_z \\ &= j\omega\mu \mathbf{k} \times \nabla H_z\end{aligned}$$

Using (3.58), the above equation is shown to be equivalent to

$$\nabla \times \nabla \times \mathbf{E}_t - \omega^2 \epsilon \mu \mathbf{E}_t = -j\omega\mu \gamma \mathbf{k} \times \mathbf{H}_t \quad (3.63)$$

Subtracting ∇ (3.61) from (3.63), we have

$$\nabla \times \nabla \times \mathbf{E}_t - \nabla \nabla \cdot \mathbf{E}_t - \omega^2 \epsilon \mu \mathbf{E}_t = -j\omega\mu \gamma \mathbf{k} \times \mathbf{H}_t - \gamma \nabla E_z \quad (3.64)$$

However, $\gamma \mathbf{k} \times$ (3.60) gives

$$\gamma^2 \mathbf{k} \times \mathbf{k} \times \mathbf{E}_t + \gamma \mathbf{k} \times \mathbf{k} \times \nabla E_z = j\omega\mu \gamma \mathbf{k} \times \mathbf{H}_t$$

which can be rewritten as follows using (2.14) and the fact that $\mathbf{E}_t$ and ∇E_z are both perpendicular to $\mathbf{k}$.

$$-\gamma^2 \mathbf{E}_t - \gamma \nabla E_z = j\omega\mu \gamma \mathbf{k} \times \mathbf{H}_t$$

Substituting this into (3.64), the desired equation for $\mathbf{E}_t$ is obtained:

$$\nabla \times \nabla \times \mathbf{E}_t - \nabla \nabla \cdot \mathbf{E}_t - k^2 \mathbf{E}_t = 0 \quad (\text{in } S) \quad (3.65)$$

where S is the waveguide cross section and

$$k^2 = \omega^2 \epsilon \mu + \gamma^2 \quad (3.66)$$

There are two boundary conditions: On the waveguide wall, the tangential component of $\mathbf{E}_t$ must vanish, i.e., $\mathbf{n} \times \mathbf{E}_t = 0$, and, in addition, E_z has to be zero, i.e., $\nabla \cdot \mathbf{E}_t = 0$ from (3.61). Thus, the problem is reduced to an eigenvalue problem of solving (3.65) under the boundary conditions

$$\mathbf{n} \times \mathbf{E}_t = 0, \quad \nabla \cdot \mathbf{E}_t = 0 \quad (\text{on } L) \quad (3.67)$$

where L indicates the waveguide wall and $\mathbf{n}$ is the outer normal unit vector. Once $\mathbf{E}_t$ is obtained, E_z can be calculated from (3.61), H_z from (3.59), and $\mathbf{H}_t$ from (3.60).

Following the procedure employed in the previous section, let us now study the eigenvalue problem for $\mathbf{E}_t$. We shall prove that the eigenvalues are real and nonnegative (zero or positive). Multiplying

$$\nabla \times \nabla \times \mathbf{E}_{tn} - \nabla \nabla \cdot \mathbf{E}_{tn} - k_n^2 \mathbf{E}_{tn} = 0 \quad (3.68)$$

by $\mathbf{E}_{in}^*$ and integrating the result over S , we have

$$\begin{aligned} & \int \{ \mathbf{E}_{in}^* \cdot \nabla \times \nabla \times \mathbf{E}_{in} - \mathbf{E}_{in}^* \cdot \nabla \nabla \cdot \mathbf{E}_{in} - k_n^2 \mathbf{E}_{in}^* \cdot \mathbf{E}_{in} \} dS \\ &= \int \{ (\nabla \times \mathbf{E}_{in}^*) \cdot (\nabla \times \mathbf{E}_{in}) + (\nabla \cdot \mathbf{E}_{in}^*) (\nabla \cdot \mathbf{E}_{in}) - k_n^2 \mathbf{E}_{in}^* \cdot \mathbf{E}_{in} \} dS \\ &= 0 \end{aligned}$$

where the formula for integration by parts and the boundary conditions are used. From the above equation, we obtain

$$k_n^2 = \frac{\int \{ |\nabla \times \mathbf{E}_{in}|^2 + |\nabla \cdot \mathbf{E}_{in}|^2 \} dS}{\int |\mathbf{E}_{in}|^2 dS} \quad (3.69)$$

which indicates that k_n^2 is real and nonnegative. Since k_n^2 is real, the same argument used in Section 3.3 shows that all the eigenfunction $\mathbf{E}_{in}$'s can be assumed to be real without the loss of generality. In order to show the orthogonality between the $\mathbf{E}_{in}$'s, let us next multiply (3.68) by $\mathbf{E}_{im}$ and subtract this same expression with the subscripts m and n interchanged. If the result is integrated over S and the formula for integrating by parts is used together with the boundary conditions, we obtain

$$(k_n^2 - k_m^2) \int \mathbf{E}_{in} \cdot \mathbf{E}_{im} dS = 0$$

It follows from this that the orthogonality relation

$$\int \mathbf{E}_{in} \cdot \mathbf{E}_{im} dS = 0 \quad (n \neq m) \quad (3.70)$$

holds between any two eigenfunctions with different eigenvalues.

Inspection of the variational expression (3.34) and a comparison of it with (3.28) and (3.69) suggests that

$$k^2(\mathbf{E}_t) = \frac{\int \{ (\nabla \times \mathbf{E}_t)^2 + (\nabla \cdot \mathbf{E}_t)^2 \} dS - 2 \oint \mathbf{n} \times \mathbf{E}_t \cdot \nabla \times \mathbf{E}_t dl}{\int \mathbf{E}_t^2 dS} \quad (3.71)$$

may be an appropriate variational expression for the present eigenvalue

problem, where $\oint dl$ indicates the contour integral along the boundary L . To verify this, let δk^2 be a first order variation corresponding to $\delta \mathbf{E}_t$, a small variation from $\mathbf{E}_t$. When $\delta \mathbf{E}_t$ is said to be small, it means that in addition to the magnitude of $\delta \mathbf{E}_t$, the magnitude of the derivatives $\nabla \cdot \delta \mathbf{E}_t$ and $\nabla \times \delta \mathbf{E}_t$ are assumed to be small as before. Neglecting the higher order terms, we have from (3.71)

$$\begin{aligned} & k^2 \int \mathbf{E}_t^2 dS + 2k^2 \int \mathbf{E}_t \cdot \delta \mathbf{E}_t dS + \delta k^2 \int \mathbf{E}_t^2 dS \\ &= \int \{(\nabla \times \mathbf{E}_t)^2 + (\nabla \cdot \mathbf{E}_t)^2\} dS + 2 \int \{\nabla \times \mathbf{E}_t \cdot \nabla \times \delta \mathbf{E}_t + \nabla \cdot \mathbf{E}_t \nabla \cdot \delta \mathbf{E}_t\} dS \\ &\quad - 2 \oint \{\mathbf{n} \times \mathbf{E}_t \cdot \nabla \times \mathbf{E}_t + \mathbf{n} \times \delta \mathbf{E}_t \cdot \nabla \times \mathbf{E}_t + \mathbf{n} \times \mathbf{E}_t \cdot \nabla \times \delta \mathbf{E}_t\} dl \quad (3.72) \end{aligned}$$

Using (3.71) and the formula for integration by parts, (3.72) reduces to

$$\begin{aligned} \delta k^2 \int \mathbf{E}_t^2 dS &= 2 \int \delta \mathbf{E}_t \cdot (\nabla \times \nabla \times \mathbf{E}_t - \nabla \nabla \cdot \mathbf{E}_t - k^2 \mathbf{E}_t) dS \\ &\quad + 2 \oint \{(\mathbf{n} \cdot \delta \mathbf{E}_t)(\nabla \cdot \mathbf{E}_t) - \mathbf{n} \times \mathbf{E}_t \cdot \nabla \times \delta \mathbf{E}_t\} dl \quad (3.73) \end{aligned}$$

It follows from (3.73) that if $\mathbf{E}_t$ is an eigenfunction, the first order variation δk^2 corresponding to any small $\delta \mathbf{E}_t$ vanishes. Conversely, if δk^2 is equal to zero for every possible small variation $\delta \mathbf{E}_t$ from $\mathbf{E}_t$, $\mathbf{E}_t$ is an eigenfunction for the following reasons.

(i) If the differential equation is not satisfied somewhere in S , $\delta \mathbf{E}_t$ can be made parallel to $\nabla \times \nabla \times \mathbf{E}_t - \nabla \nabla \cdot \mathbf{E}_t - k^2 \mathbf{E}_t$ and both $\mathbf{n} \cdot \delta \mathbf{E}_t$ and $\nabla \times \delta \mathbf{E}_t$ can be made equal to zero on L . Why $\nabla \times \delta \mathbf{E}_t$ can be made zero when $\mathbf{n} \cdot \delta \mathbf{E}_t = 0$ will be seen if we write $\nabla \times \delta \mathbf{E}_t$ in the form $\mathbf{i}_z \{(\partial \delta E_t / \partial n) - (\partial \delta E_n / \partial t)\}$, where δE_n and δE_t are the normal and tangential components of $\delta \mathbf{E}_t$, respectively.

(ii) If the differential equation is satisfied but the boundary condition $\nabla \cdot \mathbf{E}_t = 0$ is not, then $\mathbf{n} \cdot \delta \mathbf{E}_t$ can be made to have the same sign as $\nabla \cdot \mathbf{E}_t$ along L and $\nabla \times \delta \mathbf{E}_t$ can be made to vanish on L .

(iii) If the differential equation and the boundary condition $\nabla \cdot \mathbf{E}_t = 0$ are satisfied but $\mathbf{n} \times \mathbf{E}_t = 0$ is not, then $\nabla \times \delta \mathbf{E}_t$ can be chosen so as to have the same sign as $\mathbf{n} \times \mathbf{E}_t$ on L . In every case, δk^2 takes a nonzero value. Thus, we conclude that (3.71) is indeed a variational expression for k^2 .

Once the variational expression is obtained we can derive an infinite series of eigenfunctions conceptionally as we did before. Let us consider a set of all functions which are real and which satisfy the boundary conditions. Among them, let $\mathbf{E}_{t1}$ be a function which minimizes $k^2(\mathbf{E}_t)$, then, $\mathbf{E}_{t1}$ is an

eigenfunction with the smallest eigenvalue $k_1^2 \geq 0$. Next, let $\mathbf{E}_{t2}$ be a function which minimizes $k^2(\mathbf{E}_t)$ with the additional condition that it is orthogonal to $\mathbf{E}_{t1}$. Then, $\mathbf{E}_{t2}$ is another eigenfunction with $k_2^2 \geq k_1^2$. In this way, adding orthogonality conditions one by one, a series of eigenfunctions $\mathbf{E}_{t1}, \mathbf{E}_{t2}, \mathbf{E}_{t3}, \dots$ can be obtained. The corresponding eigenvalues satisfy the relation $0 \leq k_1^2 \leq k_2^2 \leq k_3^2 \dots$. Note that the $\mathbf{E}_{tn}$'s thus obtained satisfy (3.70) even if $k_n^2 = k_m^2$ as long as $n \neq m$. A proof will be given in Appendix I to show that k_n^2 increases indefinitely with n , i.e., $\lim_{n \rightarrow \infty} k_n^2 = \infty$.

Let us next prove the completeness of the set of eigenfunctions assuming the infinite growth of the eigenvalues. To do so, we first normalize all the eigenfunctions, i.e.,

$$\int \mathbf{E}_{tn}^2 dS = 1 \quad (3.74)$$

Let $\mathbf{f}$ be an arbitrary function which satisfies the boundary conditions and has derivatives $\nabla \cdot \mathbf{f}$ and $\nabla \times \mathbf{f}$ which are square-integrable over S . We define $\mathbf{f}_N$ by

$$\mathbf{f}_N = \mathbf{f} - \sum_{n=1}^{N-1} A_n \mathbf{E}_{tn} \quad (3.75)$$

where

$$A_n = \int \mathbf{f} \cdot \mathbf{E}_{tn} dS \quad (3.76)$$

and a_N^2 by

$$a_N^2 = \int \mathbf{f}_N^2 dS \quad (3.77)$$

Using the orthogonality and normalization conditions for the $\mathbf{E}_{tn}$'s, (3.77) can be written in the form

$$a_N^2 = \int (\mathbf{f} - \sum A_n \mathbf{E}_{tn})^2 dS = \int \mathbf{f}^2 dS - \sum_{n=1}^{N-1} A_n^2 \quad (3.78)$$

From the definition of $\mathbf{f}_N$, $\mathbf{f}_N/a_N$ is orthogonal to all the $\mathbf{E}_{tn}$'s for which n is smaller than N , i.e.,

$$\int (\mathbf{f}_N/a_N) \cdot \mathbf{E}_{tn} dS = 0 \quad (n < N) \quad (3.79)$$

Since $\mathbf{E}_{tN}$ gives the smallest value k_N^2 of $k^2(\mathbf{E}_t)$ under the same orthogonality and boundary conditions as for $\mathbf{f}_N/a_N$, $k^2(\mathbf{f}_N/a_N)$ cannot be less than k_N^2 . Noting that $\mathbf{f}_N/a_N$ is normalized by definition, we calculate

$k^2(\mathbf{f}_N/a_N)$ as follows:

$$\begin{aligned}
 k^2(\mathbf{f}_N/a_N) &= a_N^{-2} \int \{(\nabla \times \mathbf{f}_N)^2 + (\nabla \cdot \mathbf{f}_N)^2\} dS \\
 &= a_N^{-2} \int \{(\nabla \times \mathbf{f} - \nabla \times \sum A_n \mathbf{E}_{tn})^2 + (\nabla \cdot \mathbf{f} - \nabla \cdot \sum A_n \mathbf{E}_{tn})^2\} dS \\
 &= a_N^{-2} \int [(\nabla \times \mathbf{f})^2 + (\nabla \cdot \mathbf{f})^2 - 2 \sum A_n \{\nabla \cdot (\mathbf{f} \times \nabla \times \mathbf{E}_{tn}) \\
 &\quad + \mathbf{f} \cdot \nabla \times \nabla \times \mathbf{E}_{tn}\} - 2 \sum A_n \{\nabla \cdot (\mathbf{f} \nabla \cdot \mathbf{E}_{tn}) - \mathbf{f} \cdot \nabla \nabla \cdot \mathbf{E}_{tn}\} \\
 &\quad + \sum \sum A_n A_m \{\nabla \cdot (\mathbf{E}_{tn} \times \nabla \times \mathbf{E}_{tm}) + \mathbf{E}_{tn} \cdot \nabla \times \nabla \times \mathbf{E}_{tm}\} \\
 &\quad + \sum \sum A_n A_m \{\nabla \cdot (\mathbf{E}_{tn} \nabla \cdot \mathbf{E}_{tm}) - \mathbf{E}_{tn} \cdot \nabla \nabla \cdot \mathbf{E}_{tm}\}] dS \\
 &= a_N^{-2} \int \{(\nabla \times \mathbf{f})^2 + (\nabla \cdot \mathbf{f})^2 - 2 \sum A_n \mathbf{f} \cdot k_n^2 \mathbf{E}_{tn} \\
 &\quad + \sum \sum A_n A_m \mathbf{E}_{tn} \cdot k_m^2 \mathbf{E}_{tm}\} dS \\
 &= a_N^{-2} \int \{(\nabla \times \mathbf{f})^2 + (\nabla \cdot \mathbf{f})^2\} dS - \frac{1}{a_N^2} \sum_{n=1}^{N-1} A_n^2 k_n^2
 \end{aligned}$$

where the boundary, orthogonality and normalization conditions are used. Noting that this has to be larger than k_N^2 and that the last term is positive, we have

$$a_N^{-2} \int \{(\nabla \times \mathbf{f})^2 + (\nabla \cdot \mathbf{f})^2\} dS \geq k_N^2$$

which is equivalent to

$$a_N^2 \leq k_N^{-2} \int \{(\nabla \times \mathbf{f})^2 + (\nabla \cdot \mathbf{f})^2\} dS \quad (3.80)$$

Since the integral is finite and k_N^2 increases indefinitely, a_N^2 approaches zero with increasing N . From the definition of a_N , we have

$$\lim_{N \rightarrow \infty} \int \left(\mathbf{f} - \sum_{n=1}^{N-1} A_n \mathbf{E}_{tn} \right)^2 dS = 0 \quad (3.81)$$

which shows that $\mathbf{f}$ can be expanded in terms of the eigenfunctions.

Let us call a function defined over a two-dimensional domain piecewise-continuous when it is continuous in the domain except along a finite number of lines each having a finite length. Let $\mathbf{F}$ be a piecewise-continuous and square-integrable function defined in S , the cross section of the waveguide. The function $\mathbf{F}$ does not have to satisfy the boundary conditions, but it can be approximated by a function $\mathbf{f}$, which satisfies the boundary conditions

and has square-integrable derivatives, in the sense that

$$\int (\mathbf{F} - \mathbf{f})^2 dS < \varepsilon/4 \quad (3.82)$$

where ε is an arbitrary small positive number. This is possible because a continuous function $\mathbf{f}$ can be constructed in such a way that $\mathbf{f}$ is equal to $\mathbf{F}$ outside $S(\varepsilon)$ while both its magnitude and direction continuously change inside $S(\varepsilon)$, where $S(\varepsilon)$ indicates the small area which completely contains the lines of discontinuity of $\mathbf{F}$. It is always possible, therefore, to satisfy (3.82) by making $S(\varepsilon)$ sufficiently small. Once, (3.82) is obtained, an argument similar to the one used in the previous section to obtain (3.46) shows that

$$\lim_{N \rightarrow \infty} \int \left(\mathbf{F} - \sum_{n=1}^{N-1} A_n \mathbf{E}_{tn} \right)^2 dS = 0$$

and the $\mathbf{E}_{tn}$'s are complete. In other words, $\mathbf{F}$ can be expanded in terms of the $\mathbf{E}_{tn}$'s:

$$\mathbf{F} = \sum_{n=1}^{\infty} A_n \mathbf{E}_{tn} \quad (3.83)$$

where

$$A_n = \int \mathbf{F} \cdot \mathbf{E}_{tn} dS \quad (3.84)$$

Let $\mathbf{F} = \sum_{n=1}^{\infty} A_n \mathbf{E}_{tn}$ and $\mathbf{G} = \sum_{n=1}^{\infty} B_n \mathbf{E}_{tn}$. If $\mathbf{F}$ is equal to $\mathbf{G}$ in the sense that $\int (\mathbf{F} - \mathbf{G})^2 dS = 0$, then $A_n = B_n$ for each n and vice versa. Furthermore,

$$\int \mathbf{F} \cdot \mathbf{G} dS = \sum_{n=1}^{\infty} A_n B_n \quad (3.85)$$

The proofs for these assertion are almost identical with those for the one-dimensional case.

Now consider $\mathbf{F} \times \mathbf{k}$. Since it satisfies the conditions for the expansion to be possible, we have

$$\mathbf{F} \times \mathbf{k} = \sum_{n=1}^{\infty} \mathbf{E}_{tn} \int \mathbf{F} \times \mathbf{k} \cdot \mathbf{E}_{tn} dS$$

After multiplying both sides by $\mathbf{k} \times$, a little manipulation yields:

$$\mathbf{F} = \sum_{n=1}^{\infty} \mathbf{k} \times \mathbf{E}_{tn} \int \mathbf{F} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS$$

This means that $\mathbf{F}$ can be expanded in terms of the $\mathbf{k} \times \mathbf{E}_{tn}$'s:

$$\mathbf{F} = \sum_{n=1}^{\infty} B_n \mathbf{k} \times \mathbf{E}_{tn} \quad (3.86)$$

where

$$B_n = \int \mathbf{F} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS \quad (3.87)$$

Either (3.83) or (3.86) can be used equally for the expansion of $\mathbf{F}$; however, if $\mathbf{n} \times \mathbf{F}$ is small on L , (3.83) is more suitable since the convergence of the series is generally better. On the other hand, if $\mathbf{n} \cdot \mathbf{F}$ is small on L , since $\mathbf{n} \cdot \mathbf{k} \times \mathbf{E}_{tn} = -\mathbf{k} \cdot \mathbf{n} \times \mathbf{E}_{tn} = 0$, $\mathbf{k} \times \mathbf{E}_{tn}$ satisfies a condition similar to that for $\mathbf{F}$ and (3.86) then becomes the first choice. Although there is no general proof for the somewhat ambiguous statement above, the situation may become plausible if we consider the following extreme case. If $\mathbf{E}_{t1}$ is expanded in terms of the $\mathbf{E}_{tn}$'s, only one term is necessary. On the other hand if the $\mathbf{k} \times \mathbf{E}_{tn}$'s are used, since the tangential component of each $\mathbf{k} \times \mathbf{E}_{tn}$ on L does not vanish, a large number of terms are necessary to resemble $\mathbf{E}_{t1}$ whose tangential component on L is zero.

When $\mathbf{F}$ is complex, the real and imaginary parts can be expanded separately and then added together to obtain the same formulas as (3.83) and (3.86).

Each eigenfunction $\mathbf{E}_{tn}$ belongs to one of the following four groups, since they exhaust all possible combinations.

- I. $\nabla \times \mathbf{E}_{tn} = 0, \quad \nabla \cdot \mathbf{E}_{tn} = 0$
- II. $\nabla \times \mathbf{E}_{tn} \neq 0, \quad \nabla \cdot \mathbf{E}_{tn} = 0$
- III. $\nabla \times \mathbf{E}_{tn} = 0, \quad \nabla \cdot \mathbf{E}_{tn} \neq 0$
- IV. $\nabla \times \mathbf{E}_{tn} \neq 0, \quad \nabla \cdot \mathbf{E}_{tn} \neq 0$

Our next task is to show the following: A complete set of eigenfunctions can be derived such that each function in the set belongs to any one of the first three groups.

Using $\mathbf{E}_{tn}$ belonging to group IV, we can define two new functions,

$$\mathbf{E}_t' = A \nabla \times \nabla \times \mathbf{E}_{tn}, \quad \mathbf{E}_t'' = B \nabla \nabla \cdot \mathbf{E}_{tn}$$

where A and B are normalizing constants for $\mathbf{E}_t'$ and $\mathbf{E}_t''$, respectively. $\mathbf{E}_{tn}$ can be rewritten in the form

$$\mathbf{E}_{tn} = k_n^{-2} \{ (1/A) \mathbf{E}_t' - (1/B) \mathbf{E}_t'' \} \quad (3.88)$$

where use is made of (3.68). Since $k_n^2 \neq 0$ from (3.69), (3.88) shows that

$\mathbf{E}_{tn}$ can be expressed as a linear combination of $\mathbf{E}_t'$ and $\mathbf{E}_t''$ in which neither $\mathbf{E}_t'$ nor $\mathbf{E}_t''$ vanishes. If $\mathbf{E}_t'$ is equal to zero, we have $\nabla \times \mathbf{E}_{tn} = 0$ from (3.88) contradicting the assumption that $\mathbf{E}_{tn}$ belongs to group IV ($\nabla \times \mathbf{E}_t''$ is equal to zero due to its form, a constant times $\nabla \times \nabla(\nabla \cdot \mathbf{E}_{tn})$). Similarly, if $\mathbf{E}_t''$ is equal to zero, $\nabla \cdot \mathbf{E}_{tn}$ becomes zero contradicting the same assumption. Substituting $\mathbf{E}_t'$ into $\nabla \times \nabla \times$ (3.68), we obtain

$$\nabla \times \nabla \times \mathbf{E}_t' - k_n^2 \mathbf{E}_t' = 0$$

Since $\nabla \cdot \mathbf{E}_t'$ is equal to zero from the definition of $\mathbf{E}_t'$, $\nabla \nabla \cdot \mathbf{E}_t'$ can be added to the left-hand side without changing the value, i.e.,

$$\nabla \times \nabla \times \mathbf{E}_t' - \nabla \nabla \cdot \mathbf{E}_t' - k_n^2 \mathbf{E}_t' = 0 \quad (\text{in } S)$$

which is exactly the same differential equation as (3.65). Next, $\mathbf{n} \times \mathbf{E}_t'$ can be calculated from (3.68) as follows.

$$\mathbf{n} \times \mathbf{E}_t' = A \mathbf{n} \times \nabla \times \nabla \times \mathbf{E}_{tn} = A \mathbf{n} \times k_n^2 \mathbf{E}_{tn} + A \mathbf{n} \times \nabla(\nabla \cdot \mathbf{E}_{tn})$$

The first term on the right-hand side is zero because of the boundary condition for $\mathbf{E}_{tn}$. Except for a constant factor, the second term can be rewritten in the form

$$\begin{aligned} \mathbf{n} \times \nabla(\nabla \cdot \mathbf{E}_{tn}) &= \mathbf{n} \times \{ \mathbf{n}(\partial/\partial n) + l(\partial/\partial l) \} (\nabla \cdot \mathbf{E}_{tn}) \\ &= \mathbf{n} \times l \{ \partial(\nabla \cdot \mathbf{E}_{tn})/\partial l \} \end{aligned}$$

However, since $\nabla \cdot \mathbf{E}_{tn}$ is equal to zero along L , the derivative with respect to l also must be equal to zero, and as a result, we have $\mathbf{n} \times \mathbf{E}_t' = 0$. Since $\nabla \cdot \mathbf{E}_t'$ is always equal to zero from the definition of $\mathbf{E}_t'$, $\mathbf{E}_t'$ satisfies

$$\mathbf{n} \times \mathbf{E}_t' = 0, \quad \nabla \cdot \mathbf{E}_t' = 0 \quad (\text{on } L)$$

which give the same boundary condition as (3.67). From these observations, we conclude that $\mathbf{E}_t'$ is an eigenfunction. Since both $\mathbf{E}_t'$ and $\mathbf{E}_{tn}$ satisfy the differential equation and the boundary conditions all of which are linear, the linear combination of $\mathbf{E}_t'$ and $\mathbf{E}_{tn}$, namely $\mathbf{E}_t''$, must satisfy the same equation and boundary conditions. Thus, $\mathbf{E}_t''$ is another eigenfunction. The orthogonality relation between $\mathbf{E}_t'$ and $\mathbf{E}_t''$ can be established as follows:

$$\begin{aligned} \int \mathbf{E}_t' \cdot \mathbf{E}_t'' dS &= AB \int (\nabla \times \nabla \times \mathbf{E}_{tn}) \cdot (\nabla \nabla \cdot \mathbf{E}_{tn}) dS \\ &= AB \int \nabla \cdot (\nabla \cdot \mathbf{E}_{tn}) (\nabla \times \nabla \times \mathbf{E}_{tn}) dS - AB \int (\nabla \cdot \mathbf{E}_{tn}) \nabla \cdot \nabla \times \nabla \times \mathbf{E}_{tn} dS \end{aligned}$$

The first term on the right-hand side can be converted to a contour integral by Gauss's theorem, and since $\nabla \cdot \mathbf{E}_{tm}$ vanishes on L , it is equal to zero. The second term is, likewise, equal to zero because $\nabla \cdot \nabla \times$ vanishes.

Now suppose that eigenfunctions are obtained successively by adding an orthogonality condition each time and we find that the n th function happens to appear in group IV for the first time. Instead of taking $\mathbf{E}_{tm}$ itself, let us take $\mathbf{E}_t'$ derived from it as the n th function and $\mathbf{E}_t''$ as the $n+1$ st function. $\mathbf{E}_t'$ and $\mathbf{E}_t''$ are orthogonal to $\mathbf{E}_{tm}$ ($m < n$). For example, the proof for $\mathbf{E}_t'$ to be orthogonal to $\mathbf{E}_{tm}$ ($m < n$) is as follows. If $\mathbf{E}_{tm}$ belongs to group I or III, $\nabla \times \nabla \times \mathbf{E}_{tm}$ is equal to zero; whereas, if $\mathbf{E}_{tm}$ belongs to group II, $\nabla \times \nabla \times \mathbf{E}_{tm} = k_m^2 \mathbf{E}_{tm}$. In either case

$$\begin{aligned} \int \mathbf{E}_{tm} \cdot \mathbf{E}_t' dS &= A \int \mathbf{E}_{tm} \cdot \nabla \times \nabla \times \mathbf{E}_{tm} dS \\ &= A \int \nabla \times \nabla \times \mathbf{E}_{tm} \cdot \mathbf{E}_{tm} dS = 0 \end{aligned}$$

where the formula for integration by parts and the boundary conditions are used. The proof for $\mathbf{E}_t''$ is similar.

The $n+2$ nd function can be obtained so as to minimize the variational expression for k^2 under the condition that this function is orthogonal to all the functions up to the $n+1$ st. In this way, without the loss of completeness, a set of orthogonal eigenfunctions can be derived such that each function belongs to any one of groups I, II, or III. Hereafter we shall assume that this has been done.

The electromagnetic waves derived from the $\mathbf{E}_{tm}$'s in group I have no longitudinal components, since $E_z = 0$ and $H_z = 0$ from (3.61) and (3.59) when $\gamma \neq 0$. These waves are called transverse electromagnetic modes or, in short, TEM modes. Similarly, for the waves derived from groups II and III, we have $E_z = 0$ and $H_z = 0$, respectively. Thus, transverse electric modes, or TE modes, are derived from group II and transverse magnetic modes, or TM modes, from group III. From this, one might be tempted to conclude that all the electromagnetic fields in a waveguide with a homogeneous medium and perfectly conducting walls can be divided into three groups, TEM, TE, and TM modes. We must recall, however, that an assumption was made at the beginning of this analysis that the fields vary exponentially with z . Because of this rather stringent restriction, the above fields may not represent the most general case. There is no guarantee that an electromagnetic field does not exist which is not expressible as a linear combination of the

waves obtained above. To obtain this guarantee, we shall proceed a little further in the next section.

Let us now consider TEM modes. Since $\nabla \times \mathbf{E}_{tn} = 0$ in S and $\mathbf{n} \times \mathbf{E}_{tn} = 0$ on L for $\mathbf{E}_{tn}$ from group I, two-dimensional Helmholtz's theorem shows that it can be written in the form

$$\mathbf{E}_{tn} = \nabla \varphi \quad (3.89)$$

Since $\nabla \cdot \mathbf{E}_{tn} = 0$ in S and $\mathbf{n} \times \mathbf{E}_{tn} = 0$ on L , φ satisfies

$$\nabla \cdot \nabla \varphi = 0 \quad (\text{in } S) \quad (3.90)$$

$$\mathbf{n} \times \nabla \varphi = 0 \quad (\text{on } L) \quad (3.91)$$

The other two conditions, $\nabla \times \mathbf{E}_{tn} = 0$ in S and $\nabla \cdot \mathbf{E}_{tn} = 0$ on L are automatically satisfied. The boundary condition (3.91) is equivalent to $\partial \varphi / \partial l = 0$, which means that φ is a constant along each connected boundary. From (3.89) and (3.90), φ can be considered as a two-dimensional potential function and $\mathbf{E}_{tn}$ as the corresponding static field.

Suppose that L is a singly connected contour, then by inspection, $\varphi = \text{constant}$ satisfies (3.90) and (3.91) which gives $\mathbf{E}_{tn} = 0$. This is, of course, what we expected since in an empty space completely enclosed by a conductor wall, no static field can exist. The proof is as follows: Since φ is a constant on L , we have

$$\begin{aligned} \int \nabla \cdot (\varphi \nabla \varphi) dS &= \int \varphi \nabla \varphi \cdot \mathbf{n} dl = \text{const} \int \mathbf{n} \cdot \nabla \varphi dl \\ &= \text{const} \int (\nabla \cdot \nabla \varphi) dS = 0 \end{aligned}$$

where Gauss's theorem and (3.90) are used. On the other hand, the left-hand side can be written in the form

$$\int \nabla \cdot (\varphi \nabla \varphi) dS = \int (\nabla \varphi \cdot \nabla \varphi + \varphi \nabla \cdot \nabla \varphi) dS = \int (\nabla \varphi)^2 dS$$

where (3.90) is again used. Combining these two equations, we conclude that $\mathbf{E}_{tn} = \nabla \varphi = 0$ in S . This means that in a waveguide with a singly connected boundary, a TEM mode cannot exist.

The above argument does not hold if there are two or more independent boundaries. In this case, φ can take a different value on each boundary and since there are $(N-1)$ independent ways of assigning potentials between N conductors, $(N-1)$ independent solutions belong to group I. For instance,

in a coaxial transmission line where N is equal to two, there is only one TEM mode.

Note that since $\nabla \times \mathbf{E}_{tn} = 0$ and $\nabla \cdot \mathbf{E}_{tn} = 0$ for $\mathbf{E}_{tn}$ from group I, Eq. (3.69) shows that $k_n^2 = 0$. On the other hand, $k_n^2 \neq 0$ for groups II and III.

Let us now consider TE modes, i.e., group II. By definition, $\nabla \times \mathbf{E}_{tn}$ is not equal to zero for $\mathbf{E}_{tn}$ from group II. Let us write this nonzero $\nabla \times \mathbf{E}_{tn}$ in the form

$$\nabla \times \mathbf{E}_{tn} = \mathbf{k} k_n H_{zn} \quad (3.92)$$

This is possible because $\mathbf{k} \times \nabla \times \mathbf{E}_{tn} = \nabla(\mathbf{k} \cdot \mathbf{E}_{tn}) - (\mathbf{k} \cdot \nabla) \mathbf{E}_{tn} = 0$ and the transverse component of $\nabla \times \mathbf{E}_{tn}$ always vanishes. The notation H_{zn} is used since it is equal to the z -component of the magnetic field of the mode except for a constant factor. Substituting (3.92) into $\nabla \times$ (3.68) and noting that $k_n \neq 0$, we have

$$\nabla \times \nabla \times \mathbf{k} H_{zn} - k_n^2 \mathbf{k} H_{zn} = 0$$

Since

$$\begin{aligned} \nabla \times \nabla \times \mathbf{k} H_{zn} &= -\nabla \times \mathbf{k} \times \nabla H_{zn} = \\ &= -\mathbf{k}(\nabla \cdot \nabla H_{zn}) + (\mathbf{k} \cdot \nabla) \nabla H_{zn} = -\mathbf{k} \nabla \cdot \nabla H_{zn} \end{aligned}$$

the above equation is equivalent to

$$\nabla^2 H_{zn} + k_n^2 H_{zn} = 0 \quad (\text{in } S) \quad (3.93)$$

The boundary condition $\nabla \cdot \mathbf{E}_{tn} = 0$ on L is automatically satisfied. On the other hand, since $\mathbf{n} \times \mathbf{E}_{tn}$ can be written in the form

$$\mathbf{n} \times \mathbf{E}_{tn} = k_n^{-2} \mathbf{n} \times \nabla \times \nabla \times \mathbf{E}_{tn} = k_n^{-1} \mathbf{n} \times \nabla \times \mathbf{k} H_{zn} = k_n^{-1} \mathbf{k}(\mathbf{n} \cdot \nabla H_{zn})$$

the boundary condition $\mathbf{n} \times \mathbf{E}_{tn} = 0$ gives

$$\mathbf{n} \cdot \nabla H_z = 0 \quad (\text{on } L) \quad (3.94)$$

Suppose H_{zn} and H_{zm} correspond to $\mathbf{E}_{tn}$ and $\mathbf{E}_{tm}$, respectively, then we have

$$\begin{aligned} \int H_{zn} H_{zm} dS &= (k_n k_m)^{-1} \int \nabla \times \mathbf{E}_{tn} \cdot \nabla \times \mathbf{E}_{tm} dS \\ &= (k_m / k_n) \int \mathbf{E}_{tn} \cdot \mathbf{E}_{tm} dS \end{aligned}$$

where the formula for integration by parts and the differential equation for $\mathbf{E}_{tm}$ are used together with the conditions $\nabla \cdot \mathbf{E}_{tm} = 0$ in S and $\mathbf{n} \times \mathbf{E}_{tm} = 0$ on L . This equation shows that the H_{zn} 's obtained from the $\mathbf{E}_{tn}$'s satisfy the orthogonality and normalization conditions, i.e., they are orthonormal to

each other. If $n \neq m$, the right-hand side vanishes, and if $n = m$, it becomes unity.

It follows from this that a set of orthonormal eigenfunctions of (3.93) and (3.94) can be derived from the $\mathbf{E}_{in}$'s belonging to group II. Conversely, a set of the eigenfunctions belonging to group II can be derived from the eigenfunctions of (3.93) and (3.94) using

$$-\mathbf{k} \times \nabla H_{zn} = k_n \mathbf{E}_{in} \quad (3.95)$$

where $k_n \neq 0$ is assumed. To show this, let us first consider $\nabla \cdot \mathbf{E}_{in}$. Since

$$\nabla \cdot \mathbf{E}_{in} = -k_n^{-1} \nabla \cdot (\mathbf{k} \times \nabla H_{zn}) = k_n^{-1} \mathbf{k} \cdot (\nabla \times \nabla H_{zn}) = 0$$

always $\nabla \cdot \mathbf{E}_{in}$ vanishes. Next, the boundary condition $\mathbf{n} \times \mathbf{E}_{in} = 0$ can be checked by the following calculation.

$$\mathbf{n} \times \mathbf{E}_{in} = -k_n^{-1} \mathbf{n} \times \mathbf{k} \times \nabla H_{zn} = -k_n^{-1} \mathbf{k} (\mathbf{n} \cdot \nabla H_{zn}) = 0$$

Finally, writing (3.95) in the form

$$\nabla H_{zn} = k_n \mathbf{k} \times \mathbf{E}_{in}$$

and substituting into ∇ (3.93), we obtain

$$\nabla \times \nabla \times \mathbf{E}_{in} - k_n^2 \mathbf{E}_{in} = 0$$

where we have used

$$\begin{aligned} -\mathbf{k} \times \nabla \times \nabla \times \mathbf{E}_{in} &= -\nabla (\mathbf{k} \cdot \nabla \times \mathbf{E}_{in}) + (\mathbf{k} \cdot \nabla) (\nabla \times \mathbf{E}_{in}) \\ &= \nabla (\nabla \cdot \mathbf{k} \times \mathbf{E}_{in}) \end{aligned}$$

together with the fact that both $\nabla \times \nabla \times \mathbf{E}_{in}$ and $\mathbf{E}_{in}$ have no $\mathbf{k}$ components. Since $\nabla \cdot \mathbf{E}_{in} = 0$, the above result shows that $\mathbf{E}_{in}$ satisfies (3.68). Thus, we have shown that $\mathbf{E}_{in}$ derived from (3.95) belongs to group II. Furthermore, since

$$\begin{aligned} \int \mathbf{E}_{in} \cdot \mathbf{E}_{im} dS &= (k_n k_m)^{-1} \int \nabla H_{zn} \cdot \nabla H_{zm} dS \\ &= (k_m/k_n) \int H_{zn} H_{zm} dS \end{aligned}$$

if the H_{zn} 's are orthonormal to each other, so are the $\mathbf{E}_{in}$'s. The above discussion shows that all the $\mathbf{E}_{in}$'s in group II can be derived from the independent solutions for (3.93) and (3.94) with the condition $k_n^2 \neq 0$ and vice versa. In other words, there is a one-to-one correspondence between the $\mathbf{E}_{in}$'s in group II and the H_{zn} 's for which $k_n^2 \neq 0$. It should be noted here

that the eigenfunctions of (3.93) and (3.94) can form a complete set of orthogonal functions. The proof for the completeness is essentially the same as before if we use the variational expression

$$k^2(H_z) = \frac{\int (\nabla H_z)^2 dS}{\int H_z^2 dS} \quad (3.96)$$

In the above one-to-one correspondence, if we take all the eigenfunctions, the only missing one is that for $k_n^2 = 0$. From (3.96), $\nabla H_z = 0$ if the eigenvalue is equal to zero; therefore, $H_z = \text{constant}$ is the missing eigenfunction. Thus, we conclude that any piecewise-continuous square-integrable scalar function defined in S can be expanded in terms of the $(\mathbf{k} \cdot \nabla \times \mathbf{E}_{in}/k_n)$'s except for the constant term.

One might fear that the particular choice of functions in group III may prevent some of the $\mathbf{E}_{in}$'s in group II from appearing in the selection of the complete set of eigenfunctions thus making the $(\mathbf{k} \cdot \nabla \times \mathbf{E}_{in}/k_n)$'s incomplete even after a constant term is added. However, this kind of interference among different groups does not exist since every function in group III is automatically orthogonal to any possible function in group II. This is shown by the following calculation.

$$\begin{aligned} \int \mathbf{E}_{in} \cdot \mathbf{E}_{im} dS &= -k_n^{-2} k_m^{-2} \int (\nabla \nabla \cdot \mathbf{E}_{in}) \cdot (\nabla \times \nabla \times \mathbf{E}_{im}) dS \\ &= -k_n^{-2} k_m^{-2} \int \{ \nabla \cdot (\nabla \cdot \mathbf{E}_{in}) (\nabla \times \nabla \times \mathbf{E}_{im}) - \nabla \cdot \mathbf{E}_{in} \nabla \cdot \nabla \times \nabla \times \mathbf{E}_{im} \} dS \\ &= 0 \end{aligned}$$

where $\mathbf{E}_{in}$ and $\mathbf{E}_{im}$ are assumed to belong to group III and group II, respectively.

Let us now turn our attention to TM modes, i.e., group III. Defining E_{zn} through

$$\nabla \cdot \mathbf{E}_{in} = k_n E_{zn} \quad (3.97)$$

and substituting into (3.68), we have

$$\nabla \cdot \nabla E_{zn} + k_n^2 E_{zn} = 0 \quad (\text{in } S) \quad (3.98)$$

The boundary condition $\nabla \cdot \mathbf{E}_{in} = 0$ gives

$$E_{zn} = 0 \quad (\text{on } L) \quad (3.99)$$

On the other hand, $\mathbf{n} \times \mathbf{E}_{tn} = 0$ is satisfied automatically if (3.99) holds, since

$$\begin{aligned}\mathbf{n} \times \mathbf{E}_{tn} &= \mathbf{n} \times (-k_n^{-2}) \nabla \nabla \cdot \mathbf{E}_{tn} = -k_n^{-1} \mathbf{n} \times \nabla E_{zn} \\ &= -k_n^{-1} \mathbf{n} \times \{\mathbf{n}(\partial/\partial n) + \mathbf{l}(\partial/\partial l)\} E_{zn} = 0\end{aligned}\quad (3.100)$$

Furthermore, from

$$\begin{aligned}\int E_{zn} E_{zm} dS &= (k_n k_m)^{-1} \int \nabla \cdot \mathbf{E}_{tn} \nabla \cdot \mathbf{E}_{tm} dS \\ &= (k_m/k_n) \int \mathbf{E}_{tn} \cdot \mathbf{E}_{tm} dS\end{aligned}$$

it follows that if the $\mathbf{E}_{tn}$'s are orthonormal to each other, so are the E_{zn} 's. In other words, from the set of eigenfunctions belonging to group III, a set of orthonormal eigenfunctions for (3.98) and (3.99) are derived. Conversely, we can derive the $\mathbf{E}_{tn}$'s belonging to group III from the E_{zn} 's using the relation

$$-\nabla E_{zn} = k_n \mathbf{E}_{tn} \quad (3.101)$$

Obviously, $\nabla \times \mathbf{E}_{tn} = 0$ is satisfied. A substitution of (3.101) into (3.98) gives $\nabla \nabla \cdot \mathbf{E}_{tn} + k_n^2 \mathbf{E}_{tn} = 0$ which is equivalent to (3.68) since $\nabla \times \mathbf{E}_{tn} = 0$. The boundary condition $\nabla \cdot \mathbf{E}_{tn} = 0$ on L is obtainable when (3.98), (3.99), and (3.101) are combined. The other boundary condition $\mathbf{n} \times \mathbf{E}_{tn} = 0$ is obvious from (3.100). Furthermore, since

$$\begin{aligned}\int \mathbf{E}_{tn} \cdot \mathbf{E}_{tm} dS &= (k_n k_m)^{-1} \int \nabla E_{zn} \cdot \nabla E_{zm} dS \\ &= (k_m/k_n) \int E_{zn} \cdot E_{zm} dS\end{aligned}$$

the $\mathbf{E}_{tn}$'s thus obtained are orthonormal to each other provided that the E_{zn} 's are selected to be orthonormal. It follows from the above discussion that there is a one-to-one correspondence between the $\mathbf{E}_{tn}$'s in group III and the eigenfunctions for (3.98) and (3.99). These eigenfunctions can form a complete set of orthonormal functions. The variational expression necessary for the proof is given by

$$k^2(E_z) = \frac{\int (\nabla E_z)^2 dS - 2 \oint E_z (\partial E_z / \partial n) dl}{\int E_z^2 dS} \quad (3.102)$$

Thus, we conclude that an arbitrary piecewise-continuous square-integrable scalar function defined in S can be expanded in terms of the $(\mathbf{V} \cdot \mathbf{E}_{in}/k_n)^2$'s. In the above discussion, we did not consider the possibility of k_n^2 being equal to zero. If k_n^2 is equal to zero, $\nabla E_z = 0$ in S from (3.102) which means that E_z is a constant in S . However, since $E_z = 0$ on L , this constant must be zero and the constant term is not necessary to form a complete set in the present case.

The above discussion used $\mathbf{E}_t$. However, $\mathbf{H}_t$ can be used equally well to replace $\mathbf{E}_t$. The equation for $\mathbf{H}_t$ corresponding to (3.65) becomes

$$\nabla \times \nabla \times \mathbf{H}_t - \nabla \nabla \cdot \mathbf{H}_t - k^2 \mathbf{H}_t = 0 \quad (\text{in } S)$$

the boundary conditions are

$$\mathbf{n} \cdot \mathbf{H}_t = 0, \quad \mathbf{n} \times \nabla \times \mathbf{H}_t = 0 \quad (\text{on } L)$$

A variational expression for k^2 is given by

$$k^2(\mathbf{H}_t) = \frac{\int \{(\nabla \times \mathbf{H}_t)^2 + (\nabla \cdot \mathbf{H}_t)^2\} dS - 2 \oint (\mathbf{n} \cdot \mathbf{H}_t)(\nabla \cdot \mathbf{H}_t) dl}{\int \mathbf{H}_t^2 dS}$$

It follows from this that a complete set of orthonormal eigenfunctions can be formed. Furthermore, it can be shown that the $\mathbf{H}_{in}$'s can be chosen in such a way that each of them belongs to any one of the three groups which lead to TEM, TE, and TM modes. The proofs for these statements are very similar to those for the $\mathbf{E}_{in}$'s and hence we do not repeat them.

This has been a lengthy section, therefore, it may be worth summarizing the discussion. First, an exponential variation of the fields is assumed and then an equation is found which the transverse electric field must satisfy. This equation contains a constant which is selected so that the solutions satisfy appropriate boundary conditions. It is shown that there is an infinite number of independent solutions satisfying the boundary conditions. When properly chosen, they form a complete set of orthonormal functions, each of which belongs to one of three groups leading to TEM, TE, and TM modes. Let the $\mathbf{E}_{in}$'s indicate this set of solutions. Then, an arbitrary piecewise-continuous and square-integrable vector function defined over the waveguide cross section can be expanded in terms of the $\mathbf{E}_{in}$'s as well as in terms of the $(\mathbf{k} \times \mathbf{E}_{in})$'s. Furthermore, an arbitrary piecewise-continuous and square-integrable scalar function defined over the waveguide cross

section can be expanded in terms of the $(\mathbf{k} \cdot \nabla \times \mathbf{E}_{tn}/k_n)$'s and a constant term as well as in terms of the $(\nabla \cdot \mathbf{E}_{tn}/k_n)$'s where the terms for which $k_n^2 = 0$ are excluded. Note that $\nabla \times \mathbf{E}_{tn}$ and $\nabla \cdot \mathbf{E}_{tn}$ for $\mathbf{E}_{tn}$ from the TM and TE groups, respectively, are equal to zero and do not appear in the above expansions.

How to eliminate the assumption of the exponential field variation is a subject of the next section.

3.5 General Theory of Waveguides

In this section, we shall derive the most general form of an electromagnetic field in a straight lossless waveguide with a uniform cross section. To do so, we first observe that $\mathbf{E}$, $\mathbf{H}$, $\nabla \times \mathbf{E}$, and $\nabla \times \mathbf{H}$ in Maxwell's equations are all well behaved (i.e., piecewise-continuous and square-integrable) functions, each of which can be expanded in terms of an appropriate set of functions studied in the previous section. Then, by substituting their expanded forms into Maxwell's equations, all the expansion coefficients and, hence, the electric and magnetic fields will be determined. Since no a priori assumption is required on the functional forms of the fields such as the exponential variation with z , this method should give all the possible solutions of Maxwell's equations in the waveguide.

Considering the similarity of the boundary conditions, we expand the transverse component of $\mathbf{E}$ in terms of the $\mathbf{E}_{tn}$'s and the longitudinal component in terms of the $(\nabla \cdot \mathbf{E}_{tn}/k_n)$'s:

$$\mathbf{E} = \sum \mathbf{E}_{tn} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS + \sum \mathbf{k} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} \int \mathbf{k} \cdot \mathbf{E} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} dS \quad (3.103)$$

where the real and imaginary parts are expanded separately and then added together, and the summation is from $n = 1$ to $n = \infty$. Similarly, we use the $\mathbf{k} \times \mathbf{E}_{tn}$'s and the $(\nabla \times \mathbf{E}_{tn}/k_n)$'s for the expansion of $\mathbf{H}$:

$$\begin{aligned} \mathbf{H} = & \sum \mathbf{k} \times \mathbf{E}_{tn} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS \\ & + \sum \frac{\nabla \times \mathbf{E}_{tn}}{k_n} \int \mathbf{H} \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS + \frac{\mathbf{k}}{S^{1/2}} \int \mathbf{H} \cdot \frac{\mathbf{k}}{S^{1/2}} dS \end{aligned} \quad (3.104)$$

The last term corresponds to the constant term which is necessary to expand an arbitrary scalar function in terms of the $(\mathbf{k} \cdot \nabla \times \mathbf{E}_{tn}/k_n)$'s as we explained in Section 3.4. The normalizing factor is $1/S^{1/2}$.

Next, noting that $\nabla \times \mathbf{E}$ is an $\mathbf{H}$ -like function, we expand it in terms of the $\mathbf{k} \times \mathbf{E}_{tn}$'s and the $(\nabla \times \mathbf{E}_{tn}/k_n)$'s:

$$\begin{aligned} \nabla \times \mathbf{E} = & \sum \mathbf{k} \times \mathbf{E}_{tn} \int \nabla \times \mathbf{E} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS \\ & + \sum \frac{\nabla \times \mathbf{E}_{tn}}{k_n} \int \nabla \times \mathbf{E} \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS + \frac{\mathbf{k}}{S^{1/2}} \int \nabla \times \mathbf{E} \cdot \frac{\mathbf{k}}{S^{1/2}} dS \end{aligned} \quad (3.105)$$

Since we are going to substitute these expanded forms into Maxwell's equations and compare the expansion coefficients, let us transform each expansion coefficient in (3.105) into a form which can easily be compared with the corresponding coefficient in (3.104). To this end, we write $\nabla \times \mathbf{E}$ in the form

$$\begin{aligned} \nabla \times \mathbf{E} = & \left\{ \nabla_t + \mathbf{k} \frac{\partial}{\partial z} \right\} \times \{ \mathbf{E}_t(z) + \mathbf{k} E_z(z) \} \\ = & \nabla_t \times \mathbf{E}_t(z) - \mathbf{k} \times \nabla_t E_z(z) + \frac{\partial}{\partial z} \mathbf{k} \times \mathbf{E}_t(z) \end{aligned} \quad (3.106)$$

where $\mathbf{E}_t(z)$ and $\mathbf{k} E_z(z)$ indicate the transverse and longitudinal components of $\mathbf{E}$ emphasizing that they are functions of z in contrast to the $\mathbf{E}_{tn}$'s which are independent. The first term on the right-hand side of (3.106) has only the $\mathbf{k}$ -component while the remaining terms have only the transverse component. With the help of (3.106), the first expansion coefficients in (3.105) can be rewritten in the form

$$\begin{aligned} \int \nabla \times \mathbf{E} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS = & \int \left\{ \frac{\partial}{\partial z} \mathbf{k} \times \mathbf{E}_t(z) - \mathbf{k} \times \nabla_t E_z(z) \right\} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS \\ = & \frac{d}{dz} \int \mathbf{E}_t(z) \cdot \mathbf{E}_{tn} dS - \int \nabla_t E_z(z) \cdot \mathbf{E}_{tn} dS \end{aligned}$$

Using the relation

$$\nabla_t E_z(z) \cdot \mathbf{E}_{tn} = \nabla_t \cdot \{ E_z(z) \mathbf{E}_{tn} \} - E_z(z) \nabla \cdot \mathbf{E}_{tn}$$

the second integral on the right-hand side can be rewritten further in the form of a contour integral plus a surface integral. The final result is

$$\begin{aligned} \int \nabla \times \mathbf{E} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS = & \frac{d}{dz} \int \mathbf{E}_t(z) \cdot \mathbf{E}_{tn} dS \\ & + \int \mathbf{k} \cdot \mathbf{E} \nabla \cdot \mathbf{E}_{tn} dS - \oint E_z(z) \mathbf{n} \cdot \mathbf{E}_{tn} dl \end{aligned} \quad (3.107)$$

Similarly, we can rewrite the other coefficients as follows.

$$\begin{aligned} \int \nabla \times \mathbf{E} \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS &= \int \nabla_t \times \mathbf{E}_t(z) \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS \\ &= k_n \int \mathbf{E} \cdot \mathbf{E}_{tn} dS + \oint \mathbf{n} \times \mathbf{E}_t(z) \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dl \end{aligned} \quad (3.108)$$

$$\int \nabla \times \mathbf{E} \cdot \frac{\mathbf{k}}{S^{1/2}} dS = \int \nabla_t \times \mathbf{E}_t(z) \cdot \frac{\mathbf{k}}{S^{1/2}} dS = \oint \frac{\mathbf{k} \cdot \mathbf{n} \times \mathbf{E}_t(z)}{S^{1/2}} dl \quad (3.109)$$

For the field in a waveguide with perfect conductor walls, the last terms in (3.107) and (3.108) become zero because of the boundary conditions $E_z(z)=0$ and $\mathbf{n} \times \mathbf{E}_t(z)=0$ on L . Similarly the right-hand side of (3.109) vanishes.

For the expansion of $\nabla \times \mathbf{H}$, we use the $\mathbf{E}_{tn}$'s and the $(\mathbf{k} \nabla \cdot \mathbf{E}_{tn}/k_n)$'s.

$$\nabla \times \mathbf{H} = \sum \mathbf{E}_{tn} \int \nabla \times \mathbf{H} \cdot \mathbf{E}_{tn} dS + \frac{\mathbf{k} \nabla \cdot \mathbf{E}_{tn}}{k_n} \int \nabla \times \mathbf{H} \cdot \mathbf{k} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} dS \quad (3.110)$$

Using an expression for $\nabla \times \mathbf{H}$ similar to (3.106) and following the method employed to obtain (3.107), the coefficients in (3.110) are calculated to be

$$\int \nabla \times \mathbf{H} \cdot \mathbf{E}_{tn} dS = -\frac{d}{dz} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS + \int \mathbf{H} \cdot \nabla \times \mathbf{E}_{tn} dS \quad (3.111)$$

$$\int \nabla \times \mathbf{H} \cdot \mathbf{k} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} dS = \int \mathbf{H} \cdot \mathbf{k} \times k_n \mathbf{E}_{tn} dS \quad (3.112)$$

where $\mathbf{n} \times \mathbf{E}_{tn} = 0$ and $\nabla \cdot \mathbf{E}_{tn} = 0$ on L are used to eliminate contour integrals.

Substituting (3.103) and (3.110) into $\nabla \times \mathbf{H} = j\omega\epsilon\mathbf{E}$ and writing the transverse and longitudinal components separately, we have

$$\sum \mathbf{E}_{tn} \left\{ -\frac{d}{dz} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS + \int \mathbf{H} \cdot \nabla \times \mathbf{E}_{tn} dS \right\} = j\omega\epsilon \sum \mathbf{E}_{tn} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS \quad (3.113)$$

$$\sum \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} \int \mathbf{H} \cdot \mathbf{k} \times k_n \mathbf{E}_{tn} dS = j\omega\epsilon \sum \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} \int \mathbf{k} \cdot \mathbf{E} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} dS \quad (3.114)$$

where (3.111) and (3.112) are used. Similarly, we obtain two equations from $\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H}$:

$$\begin{aligned} \sum \mathbf{k} \times \mathbf{E}_{tn} \left\{ \frac{d}{dz} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS + \int \mathbf{k} \cdot \mathbf{E} \nabla \cdot \mathbf{E}_{tn} dS \right\} \\ = -j\omega\mu \sum \mathbf{k} \times \mathbf{E}_{tn} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS \end{aligned} \quad (3.115)$$

$$\begin{aligned} \sum \frac{\nabla \times \mathbf{E}_{tn}}{k_n} k_n \int \mathbf{E} \cdot \mathbf{E}_{tn} dS \\ = -j\omega\mu \left\{ \sum \frac{\nabla \times \mathbf{E}_{tn}}{k_n} \int \mathbf{H} \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS + \frac{\mathbf{k}}{S^{1/2}} \int \mathbf{H} \cdot \frac{\mathbf{k}}{S^{1/2}} dS \right\} \end{aligned} \quad (3.116)$$

Equating the coefficients of the corresponding $\mathbf{E}_{tn}$ on both sides of (3.113), we have

$$-\frac{d}{dz} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS + \int \mathbf{H} \cdot \nabla \times \mathbf{E}_{tn} dS = j\omega\epsilon \int \mathbf{E} \cdot \mathbf{E}_{tn} dS \quad (3.117)$$

Similarly, (3.114) gives

$$\int \mathbf{H} \cdot \mathbf{k} \times k_n \mathbf{E}_{tn} dS = j\omega\epsilon \int \mathbf{k} \cdot \mathbf{E} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} dS \quad (3.118)$$

provided that $\nabla \cdot \mathbf{E}_{tn}$ is not equal to zero. In much the same way, (3.115) gives

$$\frac{d}{dz} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS + \int \mathbf{k} \cdot \mathbf{E} \nabla \cdot \mathbf{E}_{tn} dS = -j\omega\mu \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS \quad (3.119)$$

and finally (3.116) gives

$$k_n \int \mathbf{E} \cdot \mathbf{E}_{tn} dS = -j\omega\mu \int \mathbf{H} \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS \quad (3.120)$$

and

$$\int \mathbf{H} \cdot \frac{\mathbf{k}}{S^{1/2}} dS = 0 \quad (3.121)$$

provided that $\nabla \times \mathbf{E}_{tn}$ is not equal to zero.

Suppose $\mathbf{E}_{tn}$ belongs to group 1 defined in Section 3.4 then (3.117) and (3.119) together with $\nabla \times \mathbf{E}_{tn} = 0$ and $\nabla \cdot \mathbf{E}_{tn} = 0$ give

$$\begin{aligned} -\frac{d}{dz} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS &= j\omega\epsilon \int \mathbf{E} \cdot \mathbf{E}_{tn} dS \\ \frac{d}{dz} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS &= -j\omega\mu \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS \end{aligned}$$

Eliminating $\int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS$ from these equations, a differential equation for $\int \mathbf{E} \cdot \mathbf{E}_{tn} dS$ is obtained:

$$\frac{d^2}{dz^2} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS + \omega^2 \epsilon \mu \int \mathbf{E} \cdot \mathbf{E}_{tn} dS = 0$$

The most general solution is given by

$$\int \mathbf{E} \cdot \mathbf{E}_{tn} dS = A_n e^{-\gamma_n z} + B_n e^{\gamma_n z}$$

where A_n and B_n are constants and γ_n is given by

$$\gamma_n = j\omega(\epsilon\mu)^{1/2}$$

Substituting this solution back into the original equations, we obtain

$$\int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS = Z_{0n}^{-1} (A_n e^{-\gamma_n z} - B_n e^{\gamma_n z})$$

where

$$Z_{0n} = (\mu/\epsilon)^{1/2}$$

If $\mathbf{E}_{tn}$ belongs to group II, since $\nabla \cdot \mathbf{E}_{tn} = 0$, then (3.119) reduces to

$$\frac{d}{dz} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS = -j\omega\mu \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS$$

Eliminating $\int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS$ and $\int \mathbf{H} \cdot \nabla \times \mathbf{E}_{tn} dS$ from (3.117), (3.120), and the above equation, we have

$$\frac{d^2}{dz^2} \int \mathbf{E} \cdot \mathbf{E}_{tn} dS + (\omega^2 \epsilon \mu - k_n^2) \int \mathbf{E} \cdot \mathbf{E}_{tn} dS = 0 \quad (3.122)$$

hence, the solution becomes

$$\int \mathbf{E} \cdot \mathbf{E}_{tn} dS = A_n e^{-\gamma_n z} + B_n e^{\gamma_n z}$$

provided that

$$\gamma_n^2 = k_n^2 - \omega^2 \epsilon \mu$$

is not equal to zero. Substituting the above solution back into the original equations, we obtain

$$\begin{aligned} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS &= Z_{0n}^{-1} (A_n e^{-\gamma_n z} - B_n e^{\gamma_n z}) \\ \int \mathbf{H} \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS &= -\frac{k_n}{Z_{0n} \gamma_n} (A_n e^{-\gamma_n z} + B_n e^{\gamma_n z}) \end{aligned}$$

where

$$Z_{0n} = j\omega\mu/\gamma_n$$

When $\gamma^2 = 0$, the corresponding solutions are

$$\begin{aligned}\int \mathbf{E} \cdot \mathbf{E}_{tn} dS &= C_n + D_n z \\ \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS &= -\frac{1}{j\omega\mu} D_n \\ \int \mathbf{H} \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS &= -\frac{k_n}{j\omega\mu} (C_n + D_n z)\end{aligned}\quad (3.123)$$

If $\mathbf{E}_{tn}$ belongs to group III, since $\nabla \times \mathbf{E}_{tn} = 0$, then (3.117) reduces to

$$-\frac{d}{dz} \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS = j\omega\epsilon \int \mathbf{E} \cdot \mathbf{E}_{tn} dS$$

An elimination of $\int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS$ and $\int \mathbf{k} \cdot \mathbf{E} \nabla \cdot \mathbf{E}_{tn} dS$ from (3.118), (3.119) and the above equation gives exactly the same equation for $\mathbf{E}_{tn}$ as (3.122). When $\gamma_n^2 = k_n^2 - \omega^2\epsilon\mu$ is not equal to zero, we have

$$\begin{aligned}\int \mathbf{E} \cdot \mathbf{E}_{tn} dS &= A_n e^{-\gamma_n z} + B_n e^{\gamma_n z} \\ \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS &= Z_0^{-1} (A_n e^{-\gamma_n z} - B_n e^{\gamma_n z}) \\ \int \mathbf{k} \cdot \mathbf{E} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} dS &= \frac{k_n}{\gamma_n} (A_n e^{-\gamma_n z} - B_n e^{\gamma_n z})\end{aligned}$$

where

$$Z_0 = \gamma_n / j\omega\epsilon$$

If $\gamma_n^2 = 0$, the above expressions are replaced by

$$\begin{aligned}\int \mathbf{E} \cdot \mathbf{E}_{tn} dS &= \frac{-1}{j\omega\epsilon} D_n \\ \int \mathbf{H} \cdot \mathbf{k} \times \mathbf{E}_{tn} dS &= C_n + D_n z \\ \int \mathbf{k} \cdot \mathbf{E} \frac{\nabla \cdot \mathbf{E}_{tn}}{k_n} dS &= \frac{k_n}{j\omega\epsilon} (C_n + D_n z)\end{aligned}\quad (3.124)$$

Thus, the expansion coefficients in (3.103) and (3.104) have been determined for all possible cases for $\omega \neq 0$. We conclude from this that if none of the γ_n 's are equal to zero, the most general expression for the electromagnetic

field in the waveguide when ω is not zero is given by

$$\mathbf{E} = \sum \left\{ \mathbf{E}_{tn} (A_n e^{-\gamma_n z} + B_n e^{\gamma_n z}) + \mathbf{k} \frac{\mathbf{V} \cdot \mathbf{E}_{tn}}{\gamma_n} (A_n e^{-\gamma_n z} - B_n e^{\gamma_n z}) \right\} \quad (3.125)$$

$$\mathbf{H} = \sum \left\{ \mathbf{k} \times \mathbf{E}_{tn} Z_{0n}^{-1} (A_n e^{-\gamma_n z} - B_n e^{\gamma_n z}) - \frac{\mathbf{V} \times \mathbf{E}_{tn}}{Z_{0n} \gamma_n} (A_n e^{-\gamma_n z} + B_n e^{\gamma_n z}) \right\} \quad (3.126)$$

where

$$\gamma_n^2 = k_n^2 - \omega^2 \epsilon \mu \quad (3.127)$$

and

$$\begin{aligned} Z_{0n} &= (\mu/\epsilon)^{1/2} & \text{TEM} \\ Z_{0n} &= j\omega\mu/\gamma_n & \text{TE} \\ Z_{0n} &= \gamma_n/j\omega\epsilon & \text{TM} \end{aligned} \quad (3.128)$$

In (3.128), Z_{0n} for groups I, II, and III are indicated by TEM, TE, and TM, respectively, for the obvious reason explained previously. If some of the γ_n 's are equal to zero, the coefficients of the corresponding terms must be replaced by either (3.123) or (3.124) depending on the particular group concerned.

In the previous section, it was shown that electromagnetic fields in a waveguide could be divided into three groups, TEM, TE, and TM modes, assuming the exponential field variation with z . However, because of this assumption, it was not clear whether or not all possible electromagnetic fields in a waveguide could be expressed as a linear combination of those modes. In this section, the exponential field variation was not assumed (indeed, nonexponential field variation appeared in the particular case of $\gamma_n = 0$), and we have demonstrated that every possible electromagnetic field in a waveguide is given by the above expressions and that no other functional forms need be considered. This is a far stronger assertion than we could have made before; however, without this guarantee, the discussions in Sections 3.7 and 3.8 would have little meaning.

Suppose the transverse components of both $\mathbf{E}$ and $\mathbf{H}$ are given on a reference plane perpendicular to the axis of the waveguide. Let the position of this reference plane be at $z = 0$. Then, from (3.125) and (3.126), we have

$$\mathbf{E}_t(0) = \sum \mathbf{E}_{tn} (A_n + B_n), \quad \mathbf{H}_t(0) = \sum \mathbf{k} \times \mathbf{E}_{tn} Z_{0n}^{-1} (A_n - B_n)$$

which give simultaneous equations for A_n and B_n ,

$$\int \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS = A_n + B_n, \quad \int \mathbf{H}_t(0) \cdot \mathbf{k} \times \mathbf{E}_{tn} dS = Z_{0n}^{-1} (A_n - B_n)$$

These equations determine A_n and B_n for each n and, hence, the electromagnetic field in the waveguide uniquely. In other words, the transverse components of $\mathbf{E}$ and $\mathbf{H}$ at a reference plane are sufficient to specify the electromagnetic field in the waveguide. This is true even when some of the γ_n 's are equal to zero. The transverse components of $\mathbf{E}$ on two different reference planes at $z = z_1$ and $z = z_2$ are also sufficient to specify the electromagnetic field provided that none of the $\gamma_n(z_2 - z_1)$'s are equal to $j m \pi$ where m is an integer including zero. If $\gamma_n(z_2 - z_1)$ is equal to $j m \pi$, the transverse components of $\mathbf{E}$ cannot be specified independently on the two planes, and the electromagnetic field in the waveguide is not determined uniquely. Many other ways exist for specifying the electromagnetic field uniquely, apart from the two described above, but we shall not discuss these here.

Let us now look at the expansion coefficients appearing in (3.125) and (3.126). They have exactly the same functional forms as those of voltage and current along a transmission line. We can, therefore, introduce a one-to-one correspondence between the electric and magnetic fields in a waveguide and the voltage $V_n(z)$ and current $I_n(z)$ on an infinite number of transmission lines, each representing one mode in the waveguide,

$$\mathbf{E} = \sum \left\{ \mathbf{E}_{tn} V_n(z) + \mathbf{k} \frac{\nabla \cdot \mathbf{E}_{tn}}{\gamma_n} Z_{0n} I_n(z) \right\} \quad (3.129)$$

$$\mathbf{H} = \sum \left\{ \mathbf{k} \times \mathbf{E}_{tn} I_n(z) - \nabla \times \mathbf{E}_{tn} \frac{V_n(z)}{Z_{0n} \gamma_n} \right\} \quad (3.130)$$

Since the transmission power in the waveguide is given by

$$\begin{aligned} P &= \operatorname{Re} \int \mathbf{k} \cdot \mathbf{E} \times \mathbf{H}^* dS = \operatorname{Re} \int \mathbf{E} \cdot (-\mathbf{k} \times \mathbf{H}^*) dS \\ &= \operatorname{Re} \int \sum \mathbf{E}_{tn} V_n(z) \cdot \sum \mathbf{E}_{tn} I_n^*(z) dS \\ &= \operatorname{Re} \sum V_n(z) I_n^*(z) \end{aligned} \quad (3.131)$$

it is equal to the sum of the transmission power on each transmission line separately.

For a fixed frequency ω , only a finite number of γ_n 's can become imaginary since $\lim_{n \rightarrow \infty} k_n^2 = \infty$. This means that all but a finite number of modes are in the cutoff region (cf. Section 3.2). Therefore, as long as we avoid sections of the waveguide where cutoff modes are excited, the waveguide can be represented by a finite number of transmission lines. The effect of the

cutoff modes is usually expressed by multiport networks which connect the transmission lines representing the propagating modes as we shall see in Chapter 5.

For each propagating mode, the wavelength in the guide is given by

$$\lambda_g = \lambda \{1 - (\lambda/\lambda_c)^2\}^{-1/2}$$

where λ is the free-space wavelength and λ_c the cutoff wavelength. The phase and group velocities are

$$v_p = v_0 \{1 - (\lambda/\lambda_c)^2\}^{-1/2}, \quad v_g = v_0 \{1 - (\lambda/\lambda_c)^2\}^{1/2}$$

where $v_0 = (\epsilon\mu)^{-1/2}$ and λ_c for TEM modes is infinite. These are the same formulas as those for the rectangular TE_{10} mode. The characteristic impedance is real and it decreases or increases with increasing frequency in the case of TE modes or TM modes, respectively. For TEM modes, the characteristic impedance is independent of the frequency. It may be worth mentioning that the above choice of voltage and current are somewhat arbitrary. As we explained in Section 3.2, $nV_n(z)$ and $I_n(z)/n$ can also be considered as the voltage and current of a transmission line whose characteristic impedance is $n^2 Z_{0n}$.

3.6 Examples of Waveguides

As the first example, let us again consider a rectangular waveguide. As we discussed in Section 3.4, no TEM modes exist in this case and, since the fields for TE and TM modes can all be derived from H_z and E_z , respectively, it is not necessary to deal with the vector differential equation (3.65) directly. The scalar eigenvalue problems for H_z and E_z should suffice.

For TE modes, the eigenvalue problem is given by

$$\begin{aligned} \frac{\partial^2 H_z}{\partial x^2} + \frac{\partial^2 H_z}{\partial y^2} + k^2 H_z &= 0 \\ \frac{\partial H_z}{\partial x} &= 0 \quad (x=0, \quad x=a) \\ \frac{\partial H_z}{\partial y} &= 0 \quad (y=0, \quad y=a) \end{aligned} \quad (3.132)$$

The solution can be found by inspection as follows

$$\begin{aligned} H_z &= A \cos(n\pi x/a) \cos(m\pi y/b) \\ k^2 &= (n\pi/a)^2 + (m\pi/b)^2 \end{aligned} \quad (3.133)$$

If $n = m = 0$, k^2 becomes zero and H_z , a constant. However, this is the constant term we discussed in Section 3.5, and no fields exist corresponding to this solution. If $a > b$, the smallest eigenvalue which gives nontrivial fields is obtained when $n = 1$ and $m = 0$. This corresponds to the TE_{10} mode discussed in Section 3.2.

We know that (3.133) satisfies (3.132); however, a question may arise whether or not the functions can form the complete set necessary for our discussion. The answer is yes and the proof is given below. Let f be an arbitrary well-behaved function defined in S ($0 \leq x \leq a$, $0 \leq y \leq b$). Since $\cos(n\pi x/a)$ ($n = 0, 1, 2, \dots$) are the eigenfunctions of $(d^2\varphi/dx^2) + k^2\varphi = 0$ with the boundary conditions $d\varphi/dx = 0$ at $x = 0$ and $x = a$, they form a complete set of orthogonal functions over the region from $x = 0$ to $x = a$. Similarly, $\cos(m\pi y/b)$ ($m = 0, 1, 2, \dots$) form a complete set. From the completeness of $\cos(n\pi x/a)$ ($n = 0, 1, 2, \dots$),

$$\int_0^a \left\{ f(x, y) - \sum_{n=0}^{N-1} A_n(y) \cos(n\pi x/a) \right\}^2 dx$$

can be made arbitrarily small for each y with sufficiently large N , where the $A_n(y)$'s are the expansion coefficients. Similarly, for each $A_n(y)$,

$$\int_0^b \left\{ A_n(y) - \sum_{m=0}^{M-1} B_{nm} \cos(m\pi y/b) \right\}^2 dy$$

can be made arbitrarily small with sufficiently large M , where the B_{nm} 's are the expansion coefficients. It follows from this that, for a given $\varepsilon > 0$, the relation

$$\begin{aligned} & \iint \left\{ f(x, y) - \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} B_{nm} \cos(n\pi x/a) \cos(m\pi y/b) \right\}^2 dx dy \\ &= \iint \left\{ f(x, y) - \sum_{n=0}^{N-1} A_n(y) \cos(n\pi x/a) + \sum_{n=0}^{N-1} A_n(y) \cos(n\pi x/a) \right. \\ &\quad \left. - \sum_{n=0}^{N-1} \sum_{m=0}^{M-1} B_{nm} \cos(n\pi x/a) \sin(m\pi y/b) \right\}^2 dx dy \\ &\leq 2 \iint \left\{ f(x, y) - \sum_{n=0}^{N-1} A_n(y) \cos(n\pi x/a) \right\}^2 dx dy \\ &\quad + 2 \sum_{n=0}^{N-1} a \int \left\{ A_n(y) - \sum_{m=0}^{M-1} B_{nm} \cos(m\pi y/b) \right\}^2 dy \\ &< \varepsilon \end{aligned}$$

can be satisfied with sufficiently large N and M , where Lemma 2 in Appendix I

is used, and the integration of the second term with respect to x is performed utilizing the orthogonal relations between cosine functions on the right-hand side of the first inequality ($a/2$ for nonzero n is replaced by a). However, since ϵ is arbitrary, the above relation shows that $\cos(n\pi x/a) \cos(m\pi y/b)$ ($n, m = 0, 1, 2, \dots$) form a complete set of orthogonal functions. This completes the proof.

For TM modes, the eigenvalue problem is given by

$$\frac{\partial^2 E_z}{\partial x^2} + \frac{\partial^2 E_z}{\partial y^2} + k^2 E_z = 0$$

$$E_z = 0 \quad (x = 0, \quad x = a, \quad y = 0, \quad y = b)$$
(3.134)

and the solutions are

$$E_z = A \sin(n\pi x/a) \sin(m\pi y/b), \quad k^2 = (n\pi/a)^2 + (m\pi/b)^2 \quad (3.135)$$

If either n or m becomes zero, $E_z = 0$ and no corresponding fields exist. Therefore, $n = m = 1$ gives the smallest eigenvalue in TM modes. Since this is larger than the eigenvalue for the TE_{10} mode, we conclude that among all possible waves the TE_{10} mode has the smallest eigenvalue and, hence, the lowest cutoff frequency, provided that $a > b$. Consequently, there is a frequency range in which all modes except the TE_{10} are in the cutoff region as we stated in Section 3.2 without proof.

At least, two modes, one TE and the other TM, share one eigenvalue k^2 when n and m are equal to or larger than 1; in this case, these modes are said to be degenerate to each other. When n modes share the same k^2 , they are n -fold degenerate. In a square waveguide of side length a , four TE and two TM modes share the eigenvalue $25\pi^2/a^2$ and, the degeneracy is sixfold.

Let us next consider a circular waveguide with radius a . Again, no TEM modes exist. For the TE modes, referring to (2.62) we have

$$\frac{1}{r} \frac{\partial}{\partial r} \left(r \frac{\partial H_z}{\partial r} \right) + \frac{1}{r^2} \frac{\partial^2 H_z}{\partial \theta^2} + k^2 H_z = 0, \quad \frac{\partial H_z}{\partial r} = 0 \quad (r = a) \quad (3.136)$$

where cylindrical coordinates are used to simplify the boundary condition. A method called the separation of variables is often found useful in the solution of equations like (3.136). We shall, therefore, assume the functional form of the solutions to be the product of a function of r alone and that of θ alone, i.e.,

$$H_z = u(r) v(\theta) \quad (3.137)$$

Substituting this into (3.136) and multiplying the result by r^2/uv , we have

$$\left(\frac{r}{u} \frac{d}{dr} r \frac{du}{dr} + k^2 r^2 \right) + \frac{1}{v} \frac{d^2 v}{d\theta^2} = 0$$

The quantity inside the bracket is a function of r alone while the remaining term is a function of θ alone. The equation implies that the first function is equal to the negative of the second function. However, a function of r cannot be equal to a function of θ , unless they are both constant. Thus, we have

$$\frac{1}{v} \frac{d^2 v}{d\theta^2} = -n^2, \quad \frac{r}{u} \frac{d}{dr} \left(r \frac{du}{dr} \right) + k^2 r^2 = n^2 \quad (3.138)$$

The first equation is equivalent to

$$\frac{d^2 v}{d\theta^2} + n^2 v = 0 \quad (3.139)$$

and the solution is given by

$$v = A \cos n\theta + B \sin n\theta$$

Since we return to the original position when θ is increased by 2π , the value of H_z and hence that of v has also returned to the original value. This requires n to be an integer or zero. The second equation in (3.138) is equivalent to

$$\frac{d^2 u}{dr^2} + \frac{1}{r} \frac{du}{dr} + \left(k^2 - \frac{n^2}{r^2} \right) u = 0 \quad (3.140)$$

which is obviously a linear second-order ordinary differential equation having two independent solutions. Just as two particular independent solutions of (3.139) are called cosine and sine functions and are indicated by $\cos n\theta$ and $\sin n\theta$, two particular independent solutions of (3.140) are called Bessel and Neumann functions and are indicated by $J_n(kr)$ and $N_n(kr)$, respectively. Figure 3.12 shows $J_n(kr)$ and $N_n(kr)$ versus kr for $n=0, 1$, and 2.

Since $N_n(kr)$ approaches minus infinity as r tends to zero, H_z becomes discontinuous in $S(r \leq a)$ and is not acceptable as an eigenfunction. Therefore, we are left with $J_n(kr)$ and the corresponding eigenfunction becomes

$$H_z = (A \cos n\theta + B \sin n\theta) J_n(kr) \quad (3.141)$$

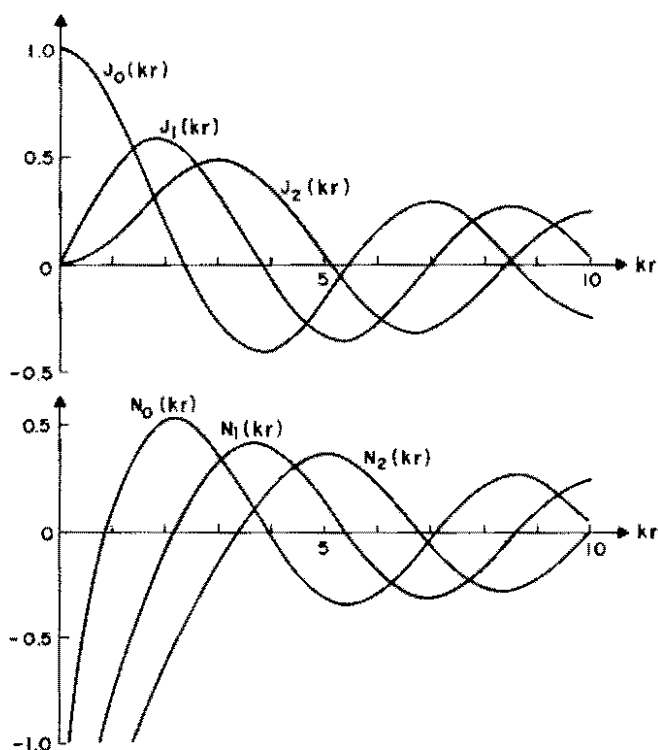


Fig. 3.12. Graphs of $J_n(kr)$ and $N_n(kr)$ versus kr .

In order to satisfy the boundary condition for H_z ,

$$\frac{dJ_n(ka)}{dr} = 0 \quad (3.142)$$

must be satisfied. Indicating the m th stationary point of $J_n(x)$ by $x = U'_{nm}$, the corresponding eigenvalue is therefore given by

$$k^2 = (U'_{nm}/a)^2 \quad (3.143)$$

Similarly, the eigenvalue problem for TM modes is given by

$$\frac{1}{r} \frac{\partial}{\partial r} \left(r \frac{\partial E_z}{\partial r} \right) + \frac{1}{r^2} \frac{\partial^2 E_z}{\partial \theta^2} + k^2 E_z = 0, \quad E_z = 0 \quad (r = a) \quad (3.144)$$

The solution for the differential equation which is continuous in $S(r \leq a)$

is given by

$$E_z = (A \cos n\theta + B \sin n\theta) J_n(kr) \quad (3.145)$$

To satisfy the boundary condition, the eigenvalue k^2 must be such that

$$J_n(ka) = 0 \quad (3.146)$$

Indicating the m th zero of $J_n(x)$ by $x = U_{nm}$, k^2 is therefore given by

$$k^2 = (U_{nm}/a)^2 \quad (3.147)$$

We have obtained an infinite number of possible solutions for both TE and TM modes. Again, however, we have to recall that an assumption was made as to the functional form of the solutions in the beginning of the analysis when we separated variables. It is, therefore, necessary to show the completeness of each set of solutions given by (3.141) and (3.145). To this end, let us first consider the orthogonality between the Bessel functions. Since U'_{np} and U'_{nq} ($p \neq q$) give different eigenvalues for (3.136), the corresponding H_z 's must be orthogonal to each other. For instance, we have

$$\int_0^a \int_0^{2\pi} \cos^2 n\theta J_n(U'_{np}r/a) J_n(U'_{nq}r/a) r \, d\theta \, dr = 0 \quad (p \neq q)$$

The integral with respect to θ does not vanish and, hence, we obtain

$$\int_0^a J_n(U'_{np}r/a) J_n(U'_{nq}r/a) r \, dr = 0 \quad (p \neq q) \quad (3.148)$$

Similarly, from the orthogonality between the E_z 's, we have

$$\int_0^a J_n(U_{np}r/a) J_n(U_{nq}r/a) r \, dr = 0 \quad (p \neq q) \quad (3.149)$$

These are the relations corresponding to (3.33) and are called the (weighted) orthogonality relations between the Bessel functions with r being the weight function.

When $p = q = m$, the left-hand sides of (3.148) and (3.149) are no longer equal to zero. Their values can be calculated as follows. Writing (3.140) in the form

$$\frac{d}{dr} \left(r \frac{du}{dr} \right) + \left(k^2 r - \frac{n^2}{r} \right) u = 0$$

and multiplying by $r(du/dr)$, we integrate the result with respect to r from

$r = 0$ to a . Integrating by parts, this leads to

$$2 \int_0^a k^2 u^2 r dr = [(r du/dr)^2]_0^a + [(k^2 r^2 - n^2) u^2]_0^a$$

Substituting $u = J_n(U'_{nm}r/a)$, $k^2 = (U'_{nm}/a)^2$ and noting that $J_n(0) = 0$ for nonzero n and $r du/dr = 0$ at $r = 0$ and $r = a$, we obtain

$$\int_0^a J_n^2(U'_{nm}r/a) r dr = \frac{1}{2} a^2 U_{nm}'^{-2} (U_{nm}'^2 - n^2) J_n^2(U'_{nm}) \quad (3.150)$$

Similarly, substituting $u = J_n(U_{nm}r/a)$, $k^2 = (U_{nm}/a)^2$ and noting that $u = 0$ at $r = 0$ and $r = a$ when $n \neq 0$, and $ru = 0$ at these points when $n = 0$, we have

$$\int_0^a J_n^2(U_{nm}r/a) r dr = \frac{1}{2} a^2 \{J_n'(U_{nm})\}^2 \quad (3.151)$$

where $J_n'(U_{nm})$ indicates the value of the derivative of the Bessel function $J_n(x)$ at $x = U_{nm}$.

We are now in a position to discuss the completeness of the eigenfunctions. For example, (3.141) with k^2 given by (3.143) can form a complete set of orthogonal functions if the relation

$$\lim_{M, N \rightarrow \infty} \int_0^a \int_0^{2\pi} \left\{ f(r, \theta) - \sum_{n=0}^N \sum_{m=1}^M (A_{nm} \cos n\theta + B_{nm} \sin n\theta) J_n(U'_{nm}r/a) \right\}^2 r d\theta dr = 0$$

holds for any well-behaved function f defined over S ($r \leq a$) with the properly chosen A_{nm} 's and B_{nm} 's. To prove that this is the case, we first observe that $\cos n\theta$ and $\sin n\theta$ ($n = 0, 1, 2, \dots$) form a complete set of orthogonal functions over the range from $\theta = 0$ to 2π as the eigenfunctions of (3.139) under the boundary conditions $v(0) = v(2\pi)$ and $v'(0) = v'(2\pi)$ (see Problem 3.10). Comparing the present problem with the proof of the completeness of the functions given in (3.133), it suffices to show that, for a given $\varepsilon > 0$, the relation

$$\int_0^a \left\{ g(r) - \sum_{m=1}^M C_{nm} J_n(U'_{nm}r/a) \right\}^2 r dr < \varepsilon$$

can be satisfied with sufficiently large M for every n , where $g(r)$ is a piecewise-continuous function with the integral $\int_0^a g^2(r) r dr$ being finite. However, this can be proved by the familiar technique using the following three relations:

the orthogonality relation (3.148), the variational expression for k^2 of (3.140) under the boundary conditions $r(du/dr) = 0$ at $r = 0$ and $r = a$,

$$k^2 = \frac{\int \{r(du/dr)^2 + (n^2/r)u^2\} dr}{\int ru^2 dr}$$

and the relation $\lim_{m \rightarrow \infty} k_{nm}^2 = \infty$ where k_{nm}^2 is the m th eigenvalue of (3.140). The infinite growth of k_{nm}^2 is obvious, since if k_{nm}^2 remained finite with increasing m , the eigenvalues of (3.136) should remain finite, but this is not the case. Similarly, the variational expression for k^2 applicable in (3.140) under the boundary condition $ru = 0$ at $r = 0$ and $r = a$ becomes

$$k^2 = \frac{\int \{r(du/dr)^2 + (n^2/r)u^2\} dr - [2ru(du/dr)]_0^a}{\int ru^2 dr}$$

and the functions given by (3.145) with k^2 satisfying (3.147) form a complete set of orthogonal functions.

Table 3.1 shows the values of U_{nm} and U'_{nm} for the first three n 's and m 's. The fields corresponding to U_{nm} and U'_{nm} are called the TM_{nm} and TE_{nm} modes, respectively. The table shows that the TE_{11} mode has the smallest eigenvalue and, hence, the lowest cutoff frequency. There are two independent TE_{11} modes corresponding to the $\cos \theta$ and $\sin \theta$ terms in (3.141); thus, the degeneracy is twofold. The TE_{01} and TM_{11} modes share the same eigenvalue, and since there are two independent TM_{11} modes, the degeneracy is threefold.

Let us now consider the TE_{11} modes. Since the 90° rotation of one field pattern around the z axis gives the other pattern, we have only to investigate

TABLE 3.1
Values of U_{nm} and U'_{nm}

n	U_{nm}			U'_{nm}		
	$m = 1$	$m = 2$	$m = 3$	$m = 1$	$m = 2$	$m = 3$
0	2.40	5.52	8.65	3.83	7.02	10.17
1	3.83	7.02	10.17	1.84	5.33	8.54
2	5.14	8.42	11.62	3.05	6.70	9.97

the term with, say, the cosine variation in (3.141). The coefficient is usually determined so as to ensure the normalization condition of $\mathbf{E}_{t1}$ derived through (3.95). However, since the normalization of H_{z1} guarantees that of $\mathbf{E}_{t1}$, let us find the normalization constant for H_{z1} . Writing the normalized function H_{z1} in the form

$$H_{z1} = K \cos \theta J_1(kr)$$

the value of K can be determined from the condition

$$\int H_{z1}^2 dS = \int_0^a \int_0^{2\pi} K^2 \cos^2 \theta J_1^2(kr) r d\theta dr = 1$$

where

$$k = (U'_{11}/a)$$

The above integral is readily performed using (3.150), and we obtain

$$K = \sqrt{2} \pi^{-1/2} \{J_1(U'_{11})\}^{-1} (U_{11}^2 - 1)^{-1/2} k$$

Using the expression for gradient in cylindrical coordinates given by (2.61), an explicit form for $\mathbf{E}_{t1}$ can be derived. Substituting the result into (3.125) and (3.126), all the components of the electric and magnetic fields are obtained.

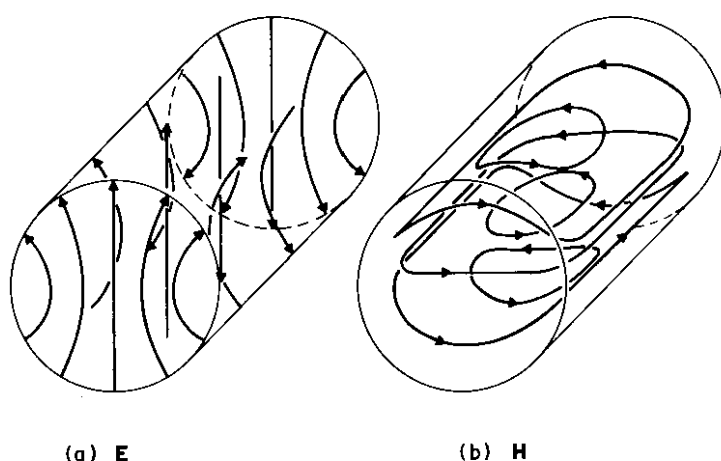
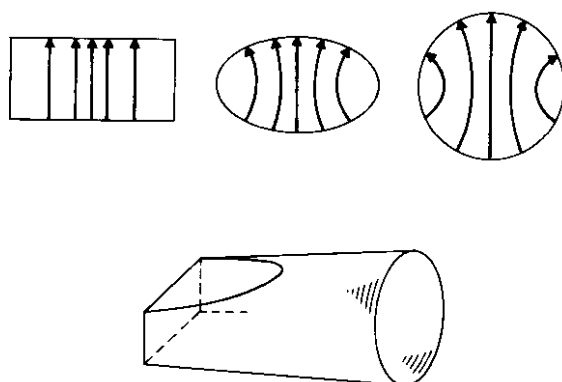
$$\begin{aligned} E_r &= K(kr)^{-1} \sin \theta J_1(kr) (Ae^{-j\beta z} + Be^{j\beta z}) \\ E_\theta &= K \cos \theta J_1'(kr) (Ae^{-j\beta z} + Be^{j\beta z}) \\ E_z &= 0 \\ H_r &= -K \cos \theta J_1'(kr) Z_0^{-1} (Ae^{-j\beta z} - Be^{j\beta z}) \\ H_\theta &= K(kr)^{-1} \sin \theta J_1(kr) Z_0^{-1} (Ae^{-j\beta z} - Be^{j\beta z}) \\ H_z &= Kk(j\omega\mu)^{-1} \cos \theta J_1(kr) (Ae^{-j\beta z} + Be^{j\beta z}) \end{aligned}$$

where A and B are constants, and

$$\beta = \{\omega^2 \epsilon \mu - (U'_{11}/a)^2\}^{1/2}, \quad Z_0 = (\omega\mu/\beta)$$

Figures 3.13(a) and (b) show the electric and magnetic field patterns, respectively, corresponding to the A terms only, i.e., those of the wave traveling in the positive z direction.

One method which will excite the circular TE_{11} mode is shown in Fig. 3.14 where a rectangular waveguide is gradually deformed into a circular waveguide so that the rectangular TE_{10} mode is gradually transformed into a circular TE_{11} mode. This is a tapered mode transducer. Since circular TE_{11} modes have the lowest cutoff frequency, they are relatively easy to handle, and because of their degeneracy, there are several interesting applications.

Fig. 3.13. Electric and magnetic fields of circular TE_{11} mode.Fig. 3.14. Transformation from rectangular TE_{10} mode to circular TE_{11} mode.

The rotary vane type attenuator is an example which has two tapered transducers to transform the circular TE_{11} mode into the rectangular TE_{10} modes at either end. A thin resistive film is placed diagonally through the axis of the circular section, as shown in Fig. 3.15(a). By rotating the circular section as a whole, one can vary the angle φ of the film with respect to the incident electric field from the rectangular section. The wave in the circular section can be considered as the sum of two TE_{11} modes, one having the electric field parallel to the film and the other normal to it. The former is attenuated by ohmic loss due to the induced current in the resistive film while the latter passes through the section without loss. Therefore, if the length of the

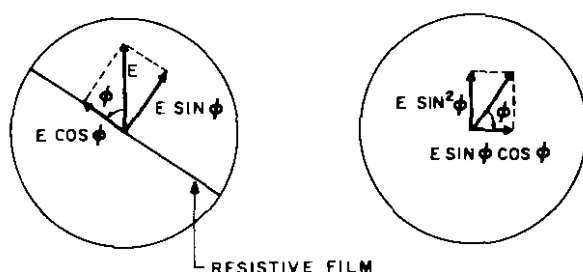


Fig. 3.15. Explanation of the principle of the rotary vane type attenuator.

circular section is sufficiently long, the first component becomes negligible whereas the second component with amplitude $E \sin \phi$ appears at the other end. This wave can in turn be considered as the sum of two modes, one with the amplitude $E \sin \phi \sin \phi$, and the other $E \sin \phi \cos \phi$ as shown in Fig. 3.15(b). If the input and output rectangular waveguides have the same orientation (not twisted to each other), the $E \sin \phi \sin \phi$ component exists through the output rectangular waveguide while the other one is reflected due to the cutoff characteristic of the rectangular TE_{01} mode. In practice, a thin resistive film is inserted in each tapered transducer section so that the cutoff mode is absorbed rather than reflected. Thus, the output electric field is proportional to $\sin^2 \phi$ and the power is proportional to $\sin^4 \phi$. When the resistive film is not ideal and has a finite thickness or is slightly misplaced, the $E \sin \phi$ component in the circular section is attenuated slightly; however, this attenuation remains the same regardless of the value of ϕ . The attenuation in the two tapered transducer sections is also constant. Hence, these three attenuations can be considered as the initial insertion loss A_0 . Therefore, as long as the $E \cos \phi$ component in the circular section is well attenuated, the total attenuation is given by

$$A = A_0 - 40 \log_{10} \sin \phi \quad (\text{dB})$$

This shows that the relative attenuation is determined solely by the angle ϕ and is independent of frequency. For this reason, rotary vane type attenuators are widely used as standards.

In the above explanation, the wave incident in the circular section was decomposed into two parts, one having substantial attenuation, and the other no attenuation. This is due to the fact that the removal of degeneracy by the resistive film takes place in such a way that each propagation constant

becomes stationary with respect to a small variation of the field. This phenomenon will be further explained in Section 3.9 for waveguides with inhomogeneous media.

3.7 Waveguide Discontinuities

In this section, let us consider an infinitely thin window across a waveguide such as the one shown in Fig. 3.16(a). The shaded area is a conducting diaphragm perpendicular to the waveguide axis and located at $z = 0$. We assume that only one propagating mode exists in the waveguide and use subscript 1 to indicate the corresponding fields. All the other modes are in the cutoff region. For convenience, let us take the reference planes of the two-port network representing the window at $z = 0$, coinciding with the dia-

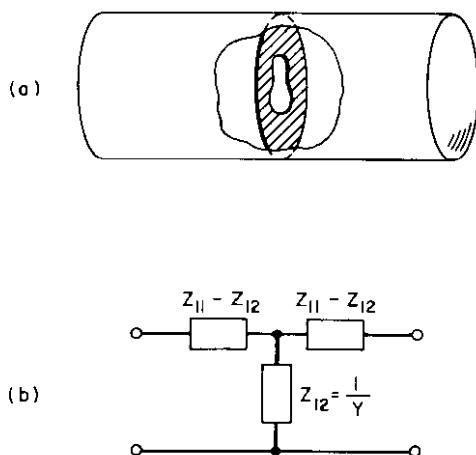


Fig. 3.16. Waveguide window and its equivalent circuit.

phragm position. The two-port network can be expressed by the equivalent circuit shown in Fig. 3.7(b). This equivalent circuit of course cannot be used unless a certain length of the transmission line representing the propagating mode is attached on each side as we explained in Section 3.2. Because of the symmetry, $Z_{11} = Z_{22}$ and the equivalent circuit reduces to Fig. 3.16(b). Now, suppose a wave with a unit magnitude is incident from $z = -\infty$; then, since this mode cannot satisfy the boundary condition imposed by the diaphragm, a number of higher modes will be generated at $z = 0$. The transverse component of the electric field is therefore expressed

in the forms

$$\mathbf{E}_t = (e^{-\gamma_1 z} + B_1 e^{\gamma_1 z}) \mathbf{E}_{t1} + \sum' B_n e^{\gamma_n z} \mathbf{E}_{tn} \quad (z < 0)$$

$$\mathbf{E}_t = \sum C_n e^{-\gamma_n z} \mathbf{E}_{tn} \quad (z > 0)$$

where $\sum'$ indicates the summation over all possible n except 1 and no wave is assumed to be reflected back at $z = \infty$. $\mathbf{E}_t$ must be continuous across the aperture and vanish on the diaphragm. Consequently, at $z = 0$ the above two expressions should give the same function $\mathbf{E}_t(0)$ which is zero on the diaphragm. From this we obtain

$$1 + B_1 = C_1 = \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \quad (3.152)$$

$$B_n = C_n = \int \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS \quad (n \geq 2) \quad (3.153)$$

The first condition shows that the voltages associated with mode 1 on each side of the window are equal. Referring to Fig. 3.16, this means that Z_{11} is equal to Z_{12} since otherwise the two voltages at the references plane cannot be equal. Thus, the window is found to be represented by a simple shunt admittance Y across the transmission line corresponding to mode 1 as we stated in Section 3.2 without proof. Furthermore, the two conditions guarantee the continuation of H_z across the aperture as well as its cancellation on the diaphragm. A comparison of (3.125) and (3.126) shows that the expansion coefficients of $\mathbf{E}_t$ are equal to the corresponding expansion coefficients of H_z and, hence, the continuity of $\mathbf{E}_t$ represented by (3.152) and (3.153) also guarantees that of H_z . To verify the cancellation of H_z , let us expand $\nabla \times \mathbf{E}_t(0)$ in terms of the $(\nabla \times \mathbf{E}_{tn}/k_n)$'s:

$$\nabla \times \mathbf{E}_t(0) = \sum \frac{\nabla \times \mathbf{E}_{tn}}{k_n} \int \nabla \times \mathbf{E}_t(0) \cdot \frac{\nabla \times \mathbf{E}_{tn}}{k_n} dS$$

Using the formula for integration by parts and the boundary condition $\mathbf{n} \times \mathbf{E}_t(0) = 0$ on the waveguide wall, the above expression reduces to

$$\nabla \times \mathbf{E}_t(0) = \sum \nabla \times \mathbf{E}_{tn} \int \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS$$

Since $\mathbf{E}_t(0)$ and hence $\nabla \times \mathbf{E}_t(0)$ vanishes on the diaphragm, the right-hand side must be equal to zero. However, using (3.126), the right-hand side is equal to $-j\omega\mu H_z$, since $Z_{0n} = j\omega\mu/\gamma_n$ for the modes with nonzero $\nabla \times \mathbf{E}_{tn}$'s, and therefore H_z vanishes on the diaphragm.

Next, let us consider $\mathbf{H}_t$ which is given by

$$\mathbf{H}_t = \mathbf{k} \times \mathbf{E}_{t1} Z_{01}^{-1} (e^{-\gamma_1 z} - B_1 e^{\gamma_1 z}) + \sum' \mathbf{k} \times \mathbf{E}_{tn} Z_{0n}^{-1} (-B_n) e^{\gamma_n z} \quad (z < 0)$$

$$\mathbf{H}_t = \mathbf{k} \times \mathbf{E}_{t1} Z_{01}^{-1} C_1 e^{-\gamma_1 z} + \sum' \mathbf{k} \times \mathbf{E}_{tn} Z_{0n}^{-1} C_n e^{-\gamma_n z} \quad (z > 0)$$

The condition that $\mathbf{H}_t$ is continuous across the aperture gives

$$\begin{aligned} \mathbf{k} \times \mathbf{E}_{t1} Z_{01}^{-1} (1 - B_1) - \sum' \mathbf{k} \times \mathbf{E}_{tn} Z_{0n}^{-1} B_n \\ = \mathbf{k} \times \mathbf{E}_{t1} Z_{01}^{-1} C_1 + \sum' \mathbf{k} \times \mathbf{E}_{tn} Z_{0n}^{-1} C_n \end{aligned}$$

Substituting (3.152) and (3.153), we have

$$\sum \mathbf{E}_{tn} Z_{0n}^{-1} B_n = 0 \quad (\text{in } S_0) \quad (3.154)$$

where S_0 indicates the aperture area. This same condition guarantees the continuity of E_z across the aperture. Referring to (3.125), E_z is continuous if the relation

$$\sum \frac{\mathbf{\nabla} \cdot \mathbf{E}_{tn}}{\gamma_n} B_n = 0 \quad (3.155)$$

holds in S_0 . Let $\mathbf{F}$ be the left-hand side of (3.154), then $\mathbf{F}$ is equal to zero in S_0 . Expanding $\mathbf{\nabla} \cdot \mathbf{F}$ in terms of the $(\mathbf{\nabla} \cdot \mathbf{E}_{tn}/k_n)$'s and using the formula for integration by parts and the boundary condition $\mathbf{\nabla} \cdot \mathbf{E}_{tn} = 0$ on the waveguide wall, we have

$$\begin{aligned} \mathbf{\nabla} \cdot \mathbf{F} &= \sum \mathbf{\nabla} \cdot \mathbf{E}_{tn} \int \mathbf{F} \cdot \mathbf{E}_{tn} dS \\ &= \sum \mathbf{\nabla} \cdot \mathbf{E}_{tn} Z_{0n}^{-1} B_n \end{aligned}$$

However, since $Z_{0n} = \gamma_n / j\omega\epsilon$, for the modes with nonzero $\mathbf{\nabla} \cdot \mathbf{E}_{tn}$'s, the left-hand side of (3.155) is found to be $j\omega\epsilon \mathbf{\nabla} \cdot \mathbf{F}$, and if we recall that $\mathbf{F}$ and hence $\mathbf{\nabla} \cdot \mathbf{F}$ vanishes in S_0 , (3.155) is proved showing the continuity of E_z across the aperture.

Since we have seen that all the boundary conditions for the fields at the window are automatically satisfied if (3.152), (3.153) and (3.154) are satisfied, let us now try to obtain the admittance Y , representing the window from these equations. Using Eqs. (3.152) and (3.153), Eq. (3.154) can be rewritten in the form

$$\mathbf{E}_{t1} = \mathbf{E}_{t1} \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS + \sum' \mathbf{E}_{tn} (Z_{01}/Z_{0n}) \int \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS \quad (\text{in } S_0) \quad (3.156)$$

This is an integral equation for $\mathbf{E}_t(0)$; therefore, if $\mathbf{E}_t(0)$ is determined, B_1

can be calculated using (3.152) and the normalized shunt admittance Y/Y_{01} of the window will be obtained from

$$1 + (Y/Y_{01}) = (1 - B_1)/(1 + B_1) \quad (3.157)$$

where use is made of (1.31) together with the knowledge that the reflection B_1 is produced by the parallel connection of the admittance Y and the line admittance Y_{01} . However, it is relatively difficult to solve the integral equation (3.156). Instead of solving (3.156), let us, therefore, express Y/Y_{01} in terms of $E_t(0)$ and rewrite it with the help of (3.156) in order to study general properties of Y/Y_{01} . From (3.152) and (3.157), Y/Y_{01} can be expressed by

$$\frac{Y}{Y_{01}} = \frac{-2B_1}{1 + B_1} = \frac{-2 \left(\int E_t(0) \cdot E_{t1} dS - 1 \right)}{\int E_t(0) \cdot E_{t1} dS} \quad (3.158)$$

Also, if we multiply both sides of (3.156) by $E_t^*(0) \cdot$ and integrating over S_0 , we have after a transposition

$$\left(1 - \int E_t(0) \cdot E_{t1} dS \right) \int E_t^*(0) \cdot E_{t1} dS = \sum' (Z_{01}/Z_{0n}) \left| \int E_t(0) \cdot E_{tn} dS \right|^2$$

Substituting this into (3.158), we obtain

$$\frac{Y}{Y_{01}} = \frac{2 \sum' (Z_{01}/Z_{0n}) \left| \int E_t(0) \cdot E_{tn} dS \right|^2}{\left| \int E_t(0) \cdot E_{t1} dS \right|^2} \quad (3.159)$$

Since all the modes except the one with subscript 1 are in the cutoff region, the (Z_{01}/Z_{0n}) 's are all pure imaginary, and hence Y/Y_{01} is a pure susceptance. Furthermore, if only one type of mode, for example, TE or TM, is excited by the window, a comparison of (3.128) and (3.159) shows that Y/Y_{01} becomes inductive or capacitive, respectively.

If we look at the expression (3.159) closely, we soon realize that the value of Y/Y_{01} remains the same when $E_t(0)$ is multiplied by a constant. This reminds us that all the variational expressions previously discussed had a similar property. Let us therefore check to determine whether or not (3.159) is a variational expression for Y/Y_{01} . To this end, we multiply both sides of (3.159) by the denominator on the right-hand side and take the variation.

The result is given by

$$\begin{aligned}
 (Y/Y_{01}) & \left\{ \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \int \Delta \mathbf{E}_t^*(0) \cdot \mathbf{E}_{t1} dS \right. \\
 & \quad \left. + \int \mathbf{E}_t^*(0) \cdot \mathbf{E}_{t1} dS \int \Delta \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \right\} + (\Delta Y/Y_{01}) \left| \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \right|^2 \\
 & = 2 \sum' (Z_{01}/Z_{0n}) \left\{ \int \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS \int \Delta \mathbf{E}_t^*(0) \cdot \mathbf{E}_{tn} dS \right. \\
 & \quad \left. + \int \mathbf{E}_t^*(0) \cdot \mathbf{E}_{tn} dS \int \Delta \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS \right\} \quad (3.160)
 \end{aligned}$$

The right-hand side can be rewritten first using (3.156) and then (3.158), as follows,

$$\begin{aligned}
 & 2 \int \Delta \mathbf{E}_t^*(0) \cdot \mathbf{E}_{t1} dS \left(1 - \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \right) \\
 & \quad - 2 \int \Delta \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \left(1 - \int \mathbf{E}_t^*(0) \cdot \mathbf{E}_{t1} dS \right) \\
 & = (Y/Y_{01}) \left\{ \int \Delta \mathbf{E}_t^*(0) \cdot \mathbf{E}_{t1} dS \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \right. \\
 & \quad \left. + \int \Delta \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \int \mathbf{E}_t^*(0) \cdot \mathbf{E}_{t1} dS \right\}
 \end{aligned}$$

where use is made of the fact that Y/Y_{01} and the (Z_{01}/Z_{0n}) 's are all pure imaginary. Substituting this into (3.160), we obtain the desired result

$$(\Delta Y/Y_{01}) \left| \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \right|^2 = 0$$

This shows that if a trial function slightly different from $\mathbf{E}_t(0)$ is substituted in (3.159), it gives an approximation at least one order higher for Y/Y_{01} than that of the trial function. Conversely, suppose that $\Delta Y/Y_{01}$ is equal to zero for all possible variations $\Delta \mathbf{E}_t(0)$ from $\mathbf{E}_t(0)$. Although this $\mathbf{E}_t(0)$ may not satisfy (3.156), it is always possible to choose a constant K such that $K\mathbf{E}_t(0)$ becomes a solution of (3.156). This will be shown as follows. If $\Delta Y/Y_{01}$ is equal to zero, (3.160) becomes

$$\begin{aligned}
 \text{Re} \int & \left\{ (Y/Y_{01}) \mathbf{E}_{t1} \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS \right. \\
 & \quad \left. - 2 \sum' (Z_{01}/Z_{0n}) \mathbf{E}_{tn} \int \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS \right\} \Delta \mathbf{E}_t^*(0) dS = 0
 \end{aligned}$$

However, since $\Delta \mathbf{E}_t^*(0)$ is arbitrary over the aperture, the terms inside the brackets must equal zero in S_0 . Equivalently, we have

$$\frac{1}{2}(Y/Y_{01}) \mathbf{E}_{t1} \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS = \sum' (Z_{01}/Z_{0n}) \mathbf{E}_{tn} \int \mathbf{E}_t(0) \cdot \mathbf{E}_{tn} dS \quad (\text{in } S_0) \quad (3.161)$$

which is linear with respect to $\mathbf{E}_t(0)$. Let K be a constant satisfying

$$K \{1 + \frac{1}{2}(Y/Y_{01})\} \int \mathbf{E}_t(0) \cdot \mathbf{E}_{t1} dS = 1$$

Replacing $\mathbf{E}_t(0)$ by $K\mathbf{E}_t(0)$ everywhere in (3.156) and substituting (3.161), we can easily see that $K\mathbf{E}_t(0)$ satisfy (3.156). We conclude from this that the solution of (3.156) is equivalent to the discovery of $\mathbf{E}_t(0)$ which makes (3.159) stationary. Thus, (3.159) is a variational expression for Y/Y_{01} in which $\mathbf{E}_t(0)$ is a solution of the integral equation (3.156). At this point the question may arise as to whether or not the solution of (3.156) is unique. The answer is generally no. Nevertheless, we can assert that Y/Y_{01} is uniquely determined. Suppose $\mathbf{A}$ and $\mathbf{B}$ are the two solutions of (3.156), then the difference $(\mathbf{A} - \mathbf{B})$ satisfies the equation

$$\mathbf{E}_{t1} \int (\mathbf{A} - \mathbf{B}) \cdot \mathbf{E}_{t1} dS + \sum' \mathbf{E}_{tn} (Z_{01}/Z_{0n}) \int (\mathbf{A} - \mathbf{B}) \cdot \mathbf{E}_{tn} dS = 0$$

Multiplying by $(\mathbf{A} - \mathbf{B})^* \cdot$ and integrating the result over S_0 , we have

$$\left| \int (\mathbf{A} - \mathbf{B}) \cdot \mathbf{E}_{t1} dS \right|^2 + \sum' (Z_{01}/Z_{0n}) \left| \int (\mathbf{A} - \mathbf{B}) \cdot \mathbf{E}_{tn} dS \right|^2 = 0$$

Since the first term is pure real while the second term is pure imaginary, the first term must be zero which implies

$$\int \mathbf{A} \cdot \mathbf{E}_{t1} dS = \int \mathbf{B} \cdot \mathbf{E}_{t1} dS$$

Consequently, the values of Y/Y_{01} for $\mathbf{E}_t(0) = \mathbf{A}$ and $\mathbf{E}_t(0) = \mathbf{B}$ must be the same as we can easily see from (3.158). The field corresponding to $\mathbf{A} - \mathbf{B}$ does not couple to the outside, and there is no way of exciting it; however, such fields can exist mathematically at certain discrete frequencies determined by the shape and size of the aperture. In practice, by deforming the aperture slightly, we can couple the fields to the outside and observe resonances in the vicinity of the discrete frequencies, the effect appearing in Y/Y_{01} for the deformed window.

Once we have established that the window can be represented by a simple shunt susceptance across the transmission line representing the dominant

mode, the value of the susceptance is best determined experimentally using a standing wave detector as described in Section 3.2, especially when the shape of the aperture is irregular.

To illustrate how to use the variational expression (3.159), let us next consider a symmetrical inductive window in a rectangular waveguide, as shown in Fig. 3.6 where the aperture width is indicated by c . The dominant mode is the rectangular TE_{10} mode, and it is obvious from the configuration of the window that only the TE_{n0} modes are necessary to satisfy the boundary condition on the diaphragm. E_{tn} and Z_{01}/Z_{0n} for these modes are given by

$$E_{tn} = \mathbf{i}_y \sqrt{2} (ab)^{-1/2} \sin(n\pi x/a)$$

$$Z_{01}/Z_{0n} = \gamma_n/\gamma_1 = -j(\lambda_g/\lambda) \{ (n\lambda/2a)^2 - 1 \}^{1/2}$$

Substituting into (3.159), we have

$$\frac{Y}{Y_{01}} = -2j \frac{\lambda_g}{\lambda} \sum' \left\{ \left(\frac{n\lambda}{2a} \right)^2 - 1 \right\}^{1/2} \frac{\left| \int E_t(0) \cdot \mathbf{i}_y \sin(n\pi x/a) dx \right|^2}{\left| \int E_t(0) \cdot \mathbf{i}_y \sin(\pi x/a) dx \right|^2}$$

Since $E_t(0)$ has only the y -component and it vanishes at the diaphragm,

$$E_t(0) = K \mathbf{i}_y \sin \left[\pi \left\{ x - \frac{1}{2}(a - c) \right\} / c \right]$$

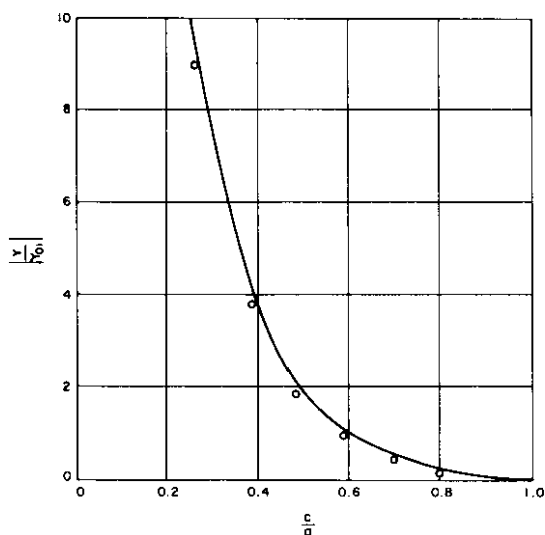


Fig. 3.17. The normalized admittance of a symmetrical inductive window ($a = 4.755$ cm; $b = 2.215$ cm; c : opening; $\lambda = 6.32$ cm).

is chosen as a trial function for $E_z(0)$ over the aperture. This is exact in the extreme case of no diaphragm. With the trial function, Y/Y_{01} becomes:

$$\frac{Y}{Y_{01}} = -2j \frac{\lambda_g \{1 - (c/a)^2\}^2}{\lambda \cos^2(\pi c/2a)} \sum_{n=3, 5, 7, \dots} \frac{\{(n\lambda/2a)^2 - 1\}^{1/2}}{\{1 - (nc/a)^2\}^2} \cos^2(n\pi c/2a)$$

Although this is an approximate expression for Y/Y_{01} , it is expected to be a good one. If $a - c$ is small, the series converges rapidly and only the first few terms are necessary to obtain as much accuracy as the trial function will permit. As an example, $|Y/Y_{01}|$ is calculated using the first three terms and shown by the solid line in Fig. 3.17 ($a = 4.755$ cm, $\lambda = 6.32$ cm). The white circles indicate the experimental results. The agreement is fairly good, considering the crude approximation of the trial function.

3.8 Effect of Wall Losses

So far we have discussed waveguides with perfectly conducting walls. In practice, however, waveguides have finite conductivity and electromagnetic energy is converted into heat due to ohmic losses when wall currents flow. Consequently, attenuation of wave amplitude takes place whenever a wave propagates down the waveguide. If attenuation is neglected in a long waveguide, an erroneous conclusion may be reached; hence, a new waveguide model is required which gives a better representation than the one with perfectly conducting walls. To obtain such a model, let us first consider the boundary conditions at the conducting walls. Let rectangular coordinates ξ, η, ζ be such that the $\xi\eta$ plane coincides with the conductor surface, and the positive direction of the ζ axis is toward the inside of the conductor. The conductor surface may be curved, but we are only interested in a very small area of the surface which can be well represented by the $\xi\eta$ plane. Since a very thin conductor sheet can shield a high frequency electromagnetic field quite effectively, let us assume that the electromagnetic field attenuates rapidly with ζ , and, hence, the change of the field in both ξ and η directions can be neglected compared to its change in the ζ direction. This means that we can assume $\partial/\partial\xi = \partial/\partial\eta = 0$. Hence, the electromagnetic field in the conductor becomes that of the plane wave discussed in Section 2.3.

Of the two solutions for the plane wave, the one with a growing exponential factor has to be abandoned to comply with the above assumption of attenuation with increasing ζ . The ratio of E_ζ to H_η as well as that of E_η to $-H_\zeta$ therefore becomes the characteristic impedance Z_w of the wall conductor. This fact can be expressed by

$$\mathbf{n} \times \mathbf{E} = Z_w \mathbf{H} \quad (3.162)$$

where $\mathbf{n}$ is the normal unit vector coinciding with the ζ axis. From (2.101), Z_w is given by

$$Z_w = (1 + j) (\omega \mu_w / 2 \sigma_w)^{1/2}$$

where σ_w and μ_w are the conductivity and permeability of the wall conductor, respectively.

Since the tangential components of $\mathbf{E}$ and $\mathbf{H}$ must be continuous across the boundary surface, as we can easily see from a discussion similar to the one presented in Section 3.1, (3.162) serves as the boundary condition for the electromagnetic field inside the waveguide. It can be rewritten in terms of the transverse and longitudinal components of the electric and magnetic fields:

$$\mathbf{n} \times \mathbf{k} E_z = Z_w \mathbf{H}_t \quad (3.163)$$

$$\mathbf{n} \times \mathbf{E}_t = Z_w \mathbf{k} H_z \quad (3.164)$$

Using (3.60) and (3.61), (3.163) becomes

$$\mathbf{n} \times \mathbf{k} \gamma^{-1} \nabla \cdot \mathbf{E}_t = Z_w (j\omega\mu)^{-1} (\gamma \mathbf{k} \times \mathbf{E}_t + \mathbf{k} \times \gamma^{-1} \nabla \nabla \cdot \mathbf{E}_t)$$

Multiplying both sides of this equation by $\gamma \mathbf{k} \times$, we have

$$\mathbf{n} \nabla \cdot \mathbf{E}_t = -Z_w (j\omega\mu)^{-1} (\gamma^2 \mathbf{E}_t + \nabla \nabla \cdot \mathbf{E}_t)$$

Since the right-hand side contains γ which is yet to be determined, let us rewrite this last equation with the help of (3.65) in the form

$$\mathbf{n} \nabla \cdot \mathbf{E}_t = -Z_w (j\omega\mu)^{-1} (\nabla \times \nabla \times \mathbf{E}_t - \omega^2 \epsilon \mu \mathbf{E}_t) \quad (3.165)$$

On the other hand, using (3.59), (3.164) becomes

$$\mathbf{n} \times \mathbf{E}_t = -Z_w (j\omega\mu)^{-1} \nabla \times \mathbf{E}_t \quad (3.166)$$

These are the boundary conditions which the electromagnetic field must satisfy in the new model of the waveguide. The problem of investigating the effect of wall losses is thus reduced to that of solving (3.65) under the boundary conditions (3.165) and (3.166), which is again an eigenvalue problem. If $\mathbf{E}_t$ is obtained for a certain k^2 , the other field components can easily be derived from it as before. It is worth mentioning that the equations involved are all linear, and therefore any superposition of the electromagnetic fields thus obtained also constitutes a solution for the field in the waveguide.

When Z_w is small, as is usually the case, it might be expected that there is a solution $\mathbf{E}_{ta}$ similar to $\mathbf{E}_{ta}$, the solution of the original eigenvalue problem whose boundary conditions are given by (3.67). Thus, if we expand $\mathbf{E}_{ta}$ in terms of a complete orthonormal set $\mathbf{E}_{tn}$ ($n = 1, 2, \dots$) in the form

$$\mathbf{E}_{ta} = \mathbf{E}_{ta} + \sum_{n \neq a} C_n \mathbf{E}_{tn} \quad (3.167)$$

where

$$C_n = \int \mathbf{E}_{t\alpha} \cdot \mathbf{E}_{tn} dS \quad (3.168)$$

the C_n 's are expected to be small and of the same order of magnitude as Z_w . Since $\mathbf{E}_{t\alpha}$ is a solution of (3.65), it satisfies

$$\nabla \times \nabla \times \mathbf{E}_{t\alpha} - \nabla \nabla \cdot \mathbf{E}_{t\alpha} - k_\alpha^2 \mathbf{E}_{t\alpha} = 0$$

where k_α^2 is the corresponding eigenvalue. To determine k_α^2 as well as the C_n 's, let us multiply the above equation by $\mathbf{E}_{tn}$ and integrate over S , the waveguide cross section. Using the formula for integration by parts twice, we have

$$\begin{aligned} & \int (\mathbf{E}_{tn} \cdot \nabla \times \nabla \times \mathbf{E}_{t\alpha} - \mathbf{E}_{tn} \cdot \nabla \nabla \cdot \mathbf{E}_{t\alpha} - k_\alpha^2 \mathbf{E}_{tn} \cdot \mathbf{E}_{t\alpha}) dS \\ &= \int \mathbf{E}_{t\alpha} \cdot (\nabla \times \nabla \times \mathbf{E}_{tn} - \nabla \nabla \cdot \mathbf{E}_{tn} - k_\alpha^2 \mathbf{E}_{tn}) dS \\ &+ \int \mathbf{n} \cdot (\mathbf{E}_{t\alpha} \times \nabla \times \mathbf{E}_{tn} - \mathbf{E}_{tn} \times \nabla \times \mathbf{E}_{t\alpha} - \mathbf{E}_{tn} \nabla \cdot \mathbf{E}_{t\alpha} + \mathbf{E}_{t\alpha} \nabla \cdot \mathbf{E}_{tn}) dl \\ &= 0 \end{aligned}$$

The first and second terms in the bracket in the second surface integral can be replaced by a single term $k_n^2 \mathbf{E}_{tn}$. The second term in the contour integral is equal to $\mathbf{n} \times \mathbf{E}_{tn} \cdot \nabla \times \mathbf{E}_{t\alpha}$ and since $\mathbf{n} \times \mathbf{E}_{tn}$ is equal to zero on L , this term vanishes. Similarly, the fourth term vanishes because $\nabla \cdot \mathbf{E}_{tn}$ is equal to zero on L . The first and third terms are $\mathbf{n} \times \mathbf{E}_{t\alpha} \cdot \nabla \times \mathbf{E}_{tn}$ and $\mathbf{E}_{tn} \cdot \mathbf{n} \nabla \cdot \mathbf{E}_{t\alpha}$, respectively. Substituting (3.166) and (3.165) into these expressions, the above equation gives

$$\begin{aligned} (k_n^2 - k_\alpha^2) \int \mathbf{E}_{t\alpha} \cdot \mathbf{E}_{tn} dS &= (j\omega\mu)^{-1} \int Z_w \{ \nabla \times \mathbf{E}_{t\alpha} \cdot \nabla \times \mathbf{E}_{tn} \\ &- (\nabla \times \nabla \times \mathbf{E}_{t\alpha} - \omega^2 \epsilon \mu \mathbf{E}_{t\alpha}) \cdot \mathbf{E}_{tn} \} dl \end{aligned} \quad (3.169)$$

Let us assume for the moment that all the C_n 's are of the same order of magnitude as Z_w . Then, $\mathbf{E}_{t\alpha}$ in (3.169) can be replaced by $\mathbf{E}_{ta}$, as long as we are interested in the approximation up to the first order of Z_w . Setting $n = a$, (3.169) becomes

$$\begin{aligned} k_a^2 - k_\alpha^2 &= (j\omega\mu)^{-1} \int Z_w \{ \nabla \times \mathbf{E}_{ta} \cdot \nabla \times \mathbf{E}_{ta} \\ &- (\nabla \times \nabla \times \mathbf{E}_{ta} - \omega^2 \epsilon \mu \mathbf{E}_{ta}) \cdot \mathbf{E}_{ta} \} dl \end{aligned} \quad (3.170)$$

from which the new eigenvalue k_a^2 can be calculated.

Since integrations similar to the one which appears on the right-hand side of (3.170) will be used several times later, let us introduce a symbol $Z_w(n, m)$ defined by

$$Z_w(n, m) = (j\omega\mu)^{-1} \int Z_w \{ \nabla \times \mathbf{E}_{in} \cdot \nabla \times \mathbf{E}_{im} - (\nabla \times \nabla \times \mathbf{E}_{in} - \omega^2 \varepsilon \mu \mathbf{E}_{in}) \cdot \mathbf{E}_{im} \} dl \quad (3.171)$$

Since we have, from the definition of $\mathbf{E}_{in}$,

$$\nabla \times \nabla \times \mathbf{E}_{in} - \omega^2 \varepsilon \mu \mathbf{E}_{in} = \nabla \nabla \cdot \mathbf{E}_{in} + \gamma_n^2 \mathbf{E}_{in} = j\omega\mu\gamma_n Z_{0n}^{-1} \mathbf{E}_{in}$$

where Z_{0n} is the characteristic impedance of the n th mode given by (3.128), then it follows that $Z_w(n, m)$ can also be written in the form

$$Z_w(n, m) = \int Z_w \{ (j\omega\mu)^{-1} \nabla \times \mathbf{E}_{in} \cdot \nabla \times \mathbf{E}_{im} - \gamma_n Z_{0n}^{-1} \mathbf{E}_{in} \cdot \mathbf{E}_{im} \} dl \quad (3.172)$$

Next, setting $n \neq a$ in (3.169) and substituting the result into (3.168), we obtain C_n to the same order of approximation as before.

$$C_n = \int \mathbf{E}_{ia} \cdot \mathbf{E}_{in} dS = (k_n^2 - k_a^2)^{-1} Z_w(a, n) \quad (3.173)$$

If every k_n^2 ($n \neq a$) is sufficiently different from k_a^2 , then every C_n becomes of the same order of magnitude as Z_w , as expected, and $\mathbf{E}_{ia}$ is determined in its expanded form. However, if $\mathbf{E}_{ia}$ is degenerate, there is at least one k_n^2 equal to k_a^2 in which case the denominator on the right-hand side of (3.173) becomes zero and the corresponding C_n can no longer be considered small. This requires a reconsideration of the assumption we made in the derivation (3.170) and (3.173) from (3.169).

To facilitate the discussion of degenerate cases, we first show that

$$\{Z_w(a, b)/Z_{0b}\} = \{Z_w(b, a)/Z_{0a}\} \quad (3.174)$$

where a and b indicate the degenerate modes and hence $k_a^2 = k_b^2$ and $\gamma_a = \gamma_b$. If one of the degenerate modes is a TM mode or if both are TEM modes, the first term in the integral on the right-hand side of (3.172) disappears, and it is obvious that (3.174) holds. If both modes a and b are TE, then $Z_{0a} = Z_{0b}$ and (3.172) show that $Z_w(a, b) = Z_w(b, a)$. However, using the relation $Z_{0a} = Z_{0b}$ again, it is obvious that (3.174) holds. We conclude from these observations that the relation (3.174) holds between any two modes degenerate to each other.

For the moment let us confine ourselves to the case of twofold degeneracy.

When mode b is degenerate to mode a , C_b may become large, and the derivations of (3.170) and (3.173) become invalid since they assume that all the C_n 's are small. To avoid this difficulty, we remove the assumption that the coefficient of E_{tb} is small in (3.167) and write E_{ta} in the form

$$E_{ta} = AE_{ta} + BE_{tb} + \sum_{n \neq a, b} C_n E_{tn} \quad (3.175)$$

where the C_n 's are expected to be small. Substituting this into (3.169), setting n equal to a and to b in turn, and taking terms up to the first order of Z_w , we obtain

$$\begin{aligned} (k_a^2 - k_a^2) A &= AZ_w(a, a) + BZ_w(b, a) \\ (k_a^2 - k_a^2) B &= AZ_w(a, b) + BZ_w(b, b) \end{aligned}$$

Transposing all the terms to the left, they become

$$\begin{aligned} \{Z_w(a, a) - (k_a^2 - k_a^2)\} A + Z_w(b, a) B &= 0 \\ Z_w(a, b) A + \{Z_w(b, b) - (k_a^2 - k_a^2)\} B &= 0 \end{aligned} \quad (3.176)$$

These are simultaneous equations for A and B . For nontrivial solutions to exist, the determinant of the coefficients must vanish. This condition gives a quadratic equation from which the value of $(k_a^2 - k_a^2)$ can be determined. Substituting this into (3.176), the ratio of A and B is determined. By substituting the first two terms on the right-hand side of (3.175) for E_{ta} everywhere on the right-hand side of (3.169) and combining the result with (3.168), the coefficient C_n can be obtained for any particular set of A and B , one of which is arbitrary. Since a quadratic equation generally has two different roots, two independent solutions E_{ta} and E_{tb} are determined by this procedure. In this case, E_{ta} and E_{tb} are no longer degenerate. As in this example, there are many cases in which degeneracy disappears due to a small perturbation given to the originally degenerate system, this phenomenon is called the removal of degeneracy (cf., discussion at the end of Section 3.6).

In order to solve (3.176), it is not necessary to follow the above general procedure. Because of the relation (3.174) between the coefficients of (3.176), the simultaneous equations are satisfied, if

$$A = Z_{0a}^{1/2} \cos \theta, \quad B = Z_{0b}^{1/2} \sin \theta \quad (3.177)$$

or

$$A = -Z_{0a}^{1/2} \sin \theta, \quad B = Z_{0b}^{1/2} \cos \theta \quad (3.178)$$

where θ is determined by

$$\tan 2\theta = \frac{2(Z_{0a}/Z_{0b})^{1/2} Z_w(a, b)}{Z_w(a, a) - Z_w(b, b)} \quad (3.179)$$

This can be checked by substituting (3.177) or (3.178) into (3.176) and showing, with the help of (3.174), that the two equations give the same $(k_a^2 - k_\alpha^2)$. The two solutions $\mathbf{E}_{t\alpha}$ and $\mathbf{E}_{t\beta}$ are, therefore, given by

$$\mathbf{E}_{t\alpha} = Z_{0a}^{1/2} \cos \theta \mathbf{E}_{ta} + Z_{0b}^{1/2} \sin \theta \mathbf{E}_{tb} + \sum_{n \neq a, b} C_n' \mathbf{E}_{tn} \quad (3.180)$$

$$\mathbf{E}_{t\beta} = -Z_{0a}^{1/2} \sin \theta \mathbf{E}_{ta} + Z_{0b}^{1/2} \cos \theta \mathbf{E}_{tb} + \sum_{n \neq a, b} C_n \mathbf{E}_{tn} \quad (3.181)$$

The corresponding eigenvalues are obtained from the following relations which are the result of substituting (3.177) and (3.178) into (3.176).

$$k_a^2 - k_\alpha^2 = Z_w(a, a) + (Z_{0b}/Z_{0a})^{1/2} Z_w(b, a) \tan \theta \quad (3.182)$$

$$k_a^2 - k_\beta^2 = Z_w(b, b) - (Z_{0a}/Z_{0b})^{1/2} Z_w(a, b) \tan \theta \quad (3.183)$$

Similarly, C_n' and C_n'' can be obtained from (3.168) if $\mathbf{E}_{t\alpha}$ is replaced by the first two terms of $\mathbf{E}_{t\alpha}$ and $\mathbf{E}_{t\beta}$, respectively.

Let $\mathbf{E}_{t\alpha}^{(0)}$ and $\mathbf{E}_{t\beta}^{(0)}$ be the principal parts; i.e., the first two terms, of $\mathbf{E}_{t\alpha}$ and $\mathbf{E}_{t\beta}$, respectively. The terms $\mathbf{E}_{t\alpha}^{(0)}$ and $\mathbf{E}_{t\beta}^{(0)}$ are not necessarily orthogonal to each other in the sense of (3.70); however, since the magnetic field $\mathbf{H}_{t\beta}^{(0)}$ corresponding to $\mathbf{E}_{t\beta}^{(0)}$ is given by

$$\mathbf{H}_{t\beta}^{(0)} = \mathbf{k} \times \left(\frac{-\sin \theta}{Z_{0a}^{1/2}} \mathbf{E}_{ta} + \frac{\cos \theta}{Z_{0b}^{1/2}} \mathbf{E}_{tb} \right)$$

we have the relation

$$\int \mathbf{E}_{t\alpha}^{(0)} \times \mathbf{H}_{t\beta}^{(0)} \cdot \mathbf{k} dS = 0$$

The subscripts α and β can be interchanged without changing the result. If modes a and b are in the frequency range where propagation is possible, Z_{0a} and Z_{0b} are both positive real, and hence the above relation can be rewritten in the form

$$\int \mathbf{E}_{t\alpha}^{(0)} \times \mathbf{H}_{t\beta}^{(0)*} \cdot \mathbf{k} dS = 0$$

This indicates to a first order approximation, that the two waves corresponding to $\mathbf{E}_{t\alpha}$ and $\mathbf{E}_{t\beta}$ carry power independently of each other.

When the degeneracy is more than twofold, (3.175) has to be replaced by

$$\mathbf{E}_{t\alpha} = A\mathbf{E}_{ta} + B\mathbf{E}_{tb} + C\mathbf{E}_{tc} + \cdots + \sum C_n \mathbf{E}_{tn}$$

where modes a, b, c , etc., are degenerate to each other and the last summa-

tion is over all possible n excluding $a, b, c, \dots$. The procedure to obtain appropriate solutions is essentially the same as before. It must be realized, however, that an m th order algebraic equation has to be solved in the case of m -fold degeneracy.

Suppose that the eigenvalue k_a^2 is obtained using the above procedure, usually called the perturbation method. Then since $k^2 = \gamma^2 + \omega^2 \epsilon \mu$, we have

$$\gamma_a^2 - \gamma_\alpha^2 = k_a^2 - k_\alpha^2$$

from which γ_α can be calculated. Let $\Delta\gamma$ be the variation from γ_a due to the wall impedance Z_w ; i.e., $\Delta\gamma = \gamma_\alpha - \gamma_a$. Then $\Delta\gamma$ is given by

$$\Delta\gamma = -\frac{1}{2}\gamma_a^{-1}(k_a^2 - k_\alpha^2)$$

If subscript a indicates a nondegenerate propagating mode, then $\gamma_a = j\beta_a$, and we obtain the following formula from (3.170) and (3.171):

$$\Delta\gamma = \frac{1}{2}j\beta_a^{-1}Z_w(a, a) \quad (3.184)$$

The real part of this expression gives the attenuation constant due to the lossy walls.

As an example illustrating how $\Delta\gamma$ is calculated in practice, let us take the TE_{10} mode in a rectangular waveguide. For this mode, $\mathbf{E}_t$ is given by

$$\mathbf{E}_t = \mathbf{i}_y \sqrt{2}(ab)^{-1/2} \sin(\pi x/a)$$

Therefore, we have

$$\begin{aligned} \nabla \times \mathbf{E}_t &= \mathbf{k}(\pi/a) \sqrt{2}(ab)^{-1/2} \cos(\pi x/a) \\ \nabla \times \nabla \times \mathbf{E}_t &= \mathbf{i}_y (\pi/a)^2 \sqrt{2}(ab)^{-1/2} \sin(\pi x/a) \end{aligned}$$

Substituting these into (3.171) and then inserting the result in (3.184), $\Delta\gamma$ is calculated to be

$$\begin{aligned} \Delta\gamma &= \frac{1}{2\beta} \frac{Z_w}{\omega\mu} \int \left\{ \left(\frac{\pi}{a} \right)^2 \frac{2}{ab} \cos^2 \frac{\pi x}{a} - \frac{2}{ab} \left(\frac{\pi^2}{a^2} - \omega^2 \epsilon \mu \right) \sin^2 \frac{\pi x}{a} \right\} dl \\ &= \frac{1}{2\beta} \frac{Z_w}{\omega\mu} \left[2 \int_0^a \left\{ \frac{2}{ab} \left(\frac{\pi}{a} \right)^2 \cos^2 \frac{\pi x}{a} + \frac{2}{ab} \left(\omega^2 \epsilon \mu - \frac{\pi^2}{a^2} \right) \sin^2 \frac{\pi x}{a} \right\} dx \right. \\ &\quad \left. + 2 \int_0^b \left(\frac{\pi}{a} \right)^2 \frac{2}{ab} dy \right] \\ &= \left\{ \omega^2 \epsilon \mu - \left(\frac{\pi}{a} \right)^2 \right\}^{-1/2} \frac{Z_w}{\omega\mu} \left\{ \frac{1}{b} \omega^2 \epsilon \mu + \frac{2}{a} \left(\frac{\pi}{a} \right)^2 \right\} \end{aligned}$$

The real part of this expression gives the attenuation constant.

In general, Z_w is proportional to $\omega^{1/2}$, and for large ω , β is proportional to ω . Therefore, the attenuation constant due to the wall losses increases in a manner approximately proportional to $\omega^{1/2}$ for large ω . On the other hand, as ω decreases toward the cutoff frequency, the attenuation constant again increases, since β decreases, and at the cutoff frequency the attenuation becomes infinite. As a result, at a certain ω , the attenuation constant takes a minimum value. For certain waves whose E_t vanishes on the wall, the term proportional to ω^2 in (3.171) disappears and the attenuation constant decreases indefinitely as ω increases. The TE_{0n} modes in circular waveguides are an example. Although the circular TE_{0n} modes are degenerate to the TM_{1n} modes, the removal of the degeneracy takes place in such a way that (3.184) is directly applicable since $Z_w(a, b)$ happens to be zero when a and b indicate the TE_{0n} and TM_{1n} modes, respectively (the degeneracy between two independent TM_{1n} modes is not removed by uniform wall impedance Z_w).

Let P be the transmission power of a wave propagating in the positive z direction and P_L be the power loss per unit length of the waveguide walls. From the conservation of energy, the decrease of P per unit distance must be equal to P_L ; i.e.,

$$P_L = -(dP/dz)$$

Since the amplitude of the wave is proportional to $\exp(-\alpha z)$ where α is the attenuation constant, P is proportional to $\exp(-2\alpha z)$ and we have

$$\alpha = -\frac{1}{2}(dP/dz)/P = \frac{1}{2}(P_L/P) \quad (3.185)$$

In order to obtain α from this formula, we calculate P_L and P assuming that the transverse electric field of the wave is approximately given by $A\mathbf{E}_{ta}$, where A is a real constant. The wall loss per unit length is calculated from the integral of the Poynting vector over a unit length of the wall. Using the boundary condition (3.162), the loss is given by

$$\begin{aligned} P_L &= \text{Re} \int \mathbf{E} \times \mathbf{H}^* \cdot \mathbf{n} \, dl = \text{Re} \int Z_w \mathbf{H} \cdot \mathbf{H}^* \, dl \\ &= A^2 \text{Re} \int Z_w \left(\frac{|\mathbf{E}_{ta}|^2}{Z_{0a}^2} + \frac{1}{\omega^2 \mu^2} |\nabla \times \mathbf{E}_{ta}|^2 \right) dl \end{aligned}$$

If $\mathbf{E}_{ta}$ belongs to TEM or TM modes, $\nabla \times \mathbf{E}_{ta} = 0$, and we have

$$P_L = -A^2 \text{Re}(\gamma_a Z_{0a})^{-1} Z_w(a, a)$$

On the other hand, if $\mathbf{E}_{ta}$ belongs to a TE mode, then $Z_{0a} = j\omega\mu/\gamma_a$ and P_L

becomes

$$P_L = -A^2 \operatorname{Re}(j\omega\mu)^{-1} Z_w(a, a)$$

Since P is given by the integral of the Poynting vector over the waveguide cross section, we have

$$P = \operatorname{Re} \int \mathbf{E} \times \mathbf{H}^* \cdot \mathbf{k} dS = A^2 Z_{0a}^{-1}$$

Substituting these expressions into (3.185), the attenuation constant is given by

$$\alpha = -\operatorname{Re} \{ \frac{1}{2} \gamma_a^{-1} Z_w(a, a) \} = \operatorname{Re} \{ \frac{1}{2} j \beta_a^{-1} Z_w(a, a) \}$$

regardless of the type of the mode under consideration. This result is identical with that obtainable from (3.184). Since the present method for calculating α utilizes the relation between the transmission power and wall loss, it is called the power-loss method. A serious drawback to the method is the possibility of obtaining wrong attenuation constants for degenerate modes since the assumption that the transverse field is approximately given by $A\mathbf{E}_{ta}$ is wrong in most cases. When $\mathbf{E}_{ta}$ is degenerate, as we discussed before, there is no easy way of finding appropriate field configurations by the power-loss method in contrast to the perturbation method previously employed.

Before closing this section, let us briefly discuss the effect of lossy media. When the conductivity σ of the medium inside the waveguide is finite, $j\omega\epsilon\mathbf{E}$ on the right-hand side of (3.52) is replaced by $(\sigma + j\omega\epsilon)\mathbf{E}$. No other modification is necessary, and all the results obtained from (3.52) and (3.53) are valid if ϵ is everywhere replaced by $\epsilon + (\sigma/j\omega)$. Thus, the relation between k_a^2 and the propagation constant γ is given by

$$k_a^2 = \gamma^2 + \omega^2 \{ \epsilon + (\sigma/j\omega) \} \mu \quad (3.186)$$

from which the effect of the conductivity on the propagation constant can be calculated. Assuming that σ is small, and writing $\gamma = \gamma_a + \Delta\gamma$, $\Delta\gamma$ becomes

$$\Delta\gamma \simeq \frac{1}{2} j \gamma_a^{-1} \omega \sigma \mu = \frac{1}{2} \beta_a^{-1} \omega \sigma \mu.$$

Therefore, when σ is taken into account, the increment in γ is real and positive for a propagating mode, which means that the attenuation constant increases but the phase constant stays the same.

3.9 Waveguides with Inhomogeneous Media

In this section, we consider the case in which ϵ and μ are functions of transverse position but are independent of z . Maxwell's equations are given

by (3.52) and (3.53) and, hence, (3.57)–(3.60) are all valid without modification. On the other hand, the divergences of (3.52) and (3.53) give $\nabla \cdot \epsilon \mathbf{E} = 0$ and $\nabla \cdot \mu \mathbf{H} = 0$, respectively, from which we obtain

$$\nabla \cdot \epsilon \mathbf{E}_t = \gamma \epsilon E_z, \quad \nabla \cdot \mu \mathbf{H}_t = \gamma \mu H_z \quad (3.187)$$

in place of (3.61) and (3.62). We can eliminate $\mathbf{H}_t$, E_z , and H_z from Eqs (3.57)–(3.60) and (3.187) leaving a differential equation involving $\mathbf{E}_t$ alone.

$$\mu \nabla \times \mu^{-1} \nabla \times \mathbf{E}_t - \nabla \epsilon^{-1} \nabla \cdot \epsilon \mathbf{E}_t - (\omega^2 \epsilon \mu + \gamma^2) \mathbf{E}_t = 0 \quad (\text{in } S) \quad (3.188)$$

If both ϵ and μ are constant in the waveguide cross section S , this equation obviously reduces to (3.65). From the requirement that the tangential component of electric field must equal zero on the waveguide walls, the boundary conditions are given by

$$\mathbf{n} \times \mathbf{E}_t = 0, \quad \nabla \cdot \epsilon \mathbf{E}_t = 0 \quad (\text{on } L) \quad (3.189)$$

Thus, we have an eigenvalue problem, namely, the differential equation (3.188) with the boundary conditions (3.189) and a constant γ^2 to be determined. Since ϵ and μ are not constant over the waveguide cross section, $\omega^2 \epsilon \mu + \gamma^2$ can no longer be considered as a constant to be determined. Once the solutions of the eigenvalue problem are obtained, all the other components of the electric and magnetic fields can be calculated using Eqs. (3.57)–(3.60) and (3.187) as in the case with a uniform medium.

Let $\mathbf{E}_{tm}$ and $\mathbf{E}_{tn}$ be two different eigenfunctions and γ_m^2 and γ_n^2 be the corresponding eigenvalues. Since $\mathbf{E}_{tm}$ is a solution of (3.188), we have

$$\mu \nabla \times \mu^{-1} \nabla \times \mathbf{E}_{tm} - \nabla \epsilon^{-1} \nabla \cdot \epsilon \mathbf{E}_{tm} - (\omega^2 \epsilon \mu + \gamma_m^2) \mathbf{E}_{tm} = 0$$

Multiplying by $(\nabla \times \mu^{-1} \nabla \times \mathbf{E}_{tn} - \omega^2 \epsilon \mathbf{E}_{tn}) \cdot$ and integrating the result over S , we obtain

$$\begin{aligned} & \int \mu (\nabla \times \mu^{-1} \nabla \times \mathbf{E}_{tm} - \omega^2 \epsilon \mathbf{E}_{tm}) \cdot (\nabla \times \mu^{-1} \nabla \times \mathbf{E}_{tn} - \omega^2 \epsilon \mathbf{E}_{tn}) dS \\ & - \int \omega^2 \epsilon^{-1} (\nabla \cdot \epsilon \mathbf{E}_{tm}) (\nabla \cdot \epsilon \mathbf{E}_{tn}) dS \\ & - \gamma_m^2 \left\{ \int \mu^{-1} (\nabla \times \mathbf{E}_{tm}) \cdot (\nabla \times \mathbf{E}_{tn}) dS - \int \omega^2 \epsilon \mathbf{E}_{tm} \cdot \mathbf{E}_{tn} dS \right\} = 0 \end{aligned} \quad (3.190)$$

where use is made of (3.189). Interchanging the subscripts m and n and subtracting the result from (3.190), we have

$$(\gamma_n^2 - \gamma_m^2) \left\{ \int \mu^{-1} (\nabla \times \mathbf{E}_{tm}) \cdot (\nabla \times \mathbf{E}_{tn}) dS - \int \omega^2 \epsilon \mathbf{E}_{tm} \cdot \mathbf{E}_{tn} dS \right\} = 0$$

This shows that, between two eigenfunctions with different eigenvalues, there is an orthogonality relation expressed by

$$\int \mu^{-1} (\nabla \times \mathbf{E}_{tm}) \cdot (\nabla \times \mathbf{E}_{tn}) dS - \int \omega^2 \epsilon \mathbf{E}_{tm} \cdot \mathbf{E}_{tn} dS = 0 \quad (3.191)$$

Let $\mathbf{H}_{tn}$ be the magnetic field corresponding to $\mathbf{E}_{tn}$. From Eqs. (3.57)–(3.60), $\mathbf{k} \times \mathbf{H}_{tn}$ is expressed in terms of $\mathbf{E}_{tn}$ as follows:

$$\mathbf{k} \times \mathbf{H}_{tn} = -(j\omega\gamma_n)^{-1} (\nabla \times \mu^{-1} \nabla \times \mathbf{E}_{tn} - \omega^2 \epsilon \mathbf{E}_{tn}) \quad (3.192)$$

Therefore, the above orthogonality relation is equivalent to

$$\int \mathbf{k} \cdot \mathbf{E}_{tm} \times \mathbf{H}_{tn} dS = 0 \quad (3.193)$$

When the orthogonality relation (3.191) holds, we can derive another orthogonality relation from (3.190) given by

$$\begin{aligned} \int \mu (\nabla \times \mu^{-1} \nabla \times \mathbf{E}_{tm} - \omega^2 \epsilon \mathbf{E}_{tm}) \cdot (\nabla \times \mu^{-1} \nabla \times \mathbf{E}_{tn} - \omega^2 \epsilon \mathbf{E}_{tn}) dS \\ - \int \omega^2 \epsilon^{-1} (\nabla \cdot \epsilon \mathbf{E}_{tm}) (\nabla \cdot \epsilon \mathbf{E}_{tn}) dS = 0 \end{aligned} \quad (3.194)$$

Using (3.57), (3.187), and (3.192), this can be rewritten in the form

$$\int \epsilon^{-1} (\nabla \times \mathbf{H}_{tm}) \cdot (\nabla \times \mathbf{H}_{tn}) dS - \int \omega^2 \mu \mathbf{H}_{tm} \cdot \mathbf{H}_{tn} dS = 0 \quad (3.195)$$

which is the dual of (3.191).

For waveguides with homogeneous media, we found the simple orthogonality relation given by (3.70). In general, however, it is impossible to find a corresponding simple orthogonality relation when the medium is inhomogeneous; i.e.,

$$\int \mathbf{E}_{tm} \cdot \mathbf{E}_{tn} \neq 0, \quad \int \epsilon \mathbf{E}_{tm} \cdot \mathbf{E}_{tn} \neq 0$$

A variational expression for γ^2 is given by

$$\gamma^2(\mathbf{E}_t) = \frac{\int \mu (\nabla \times \mu^{-1} \nabla \times \mathbf{E}_t - \omega^2 \epsilon \mathbf{E}_t)^2 dS - \int \omega^2 \epsilon^{-1} (\nabla \cdot \epsilon \mathbf{E}_t)^2 dS}{\int \mu^{-1} (\nabla \times \mathbf{E}_t)^2 dS - \int \omega^2 \epsilon \mathbf{E}_t^2 dS} \quad (3.196)$$

In the derivation, an assumption is made that the denominator does not vanish. Although it is likely that the eigenfunctions of (3.188) and (3.189) form a complete set of orthonormal functions, a procedure similar to the one employed in Section 3.4 fails to prove the completeness when use is made of (3.196). To the author's knowledge, the completeness has not yet been proved. However, since (3.196) is a variational expression, it can be used advantageously to obtain approximate values for the γ^2 's.

If the waveguide medium is homogeneous, we could classify the fields into three groups: TEM, TE, and TM modes. This kind of classification, however, is not possible if the medium is inhomogeneous. For almost all modes, both E_z and H_z are finite.

It is a straightforward process to extend the discussion to include the effects of finite conductivity in the medium, we have only to replace ϵ by $\epsilon + (\sigma/j\omega)$ everywhere in the above discussion. In the rotary vane attenuator discussed in Section 3.6, the degeneracy of two TE_{11} modes is removed in such a way that the expression (3.196) with ϵ replaced by $\epsilon + (\sigma/j\omega)$ becomes stationary for each independent mode. The modes with the electric fields perpendicular and parallel to the resistive film give the smallest and the largest attenuations, respectively, and, hence, they are independent of each other. When the film resistivity is small, the field configuration for the second mode will be considerably distorted by the presence of the film. Nevertheless, the discussion for the total attenuation of the attenuator remains unchanged.

PROBLEMS

- 3.1 Using the variational expression given by (3.34) and a trial function $\sin(\pi x/a)$, calculate the first eigenvalue of the following eigenvalue problem.

$$(d^2 E_y / dx^2) + k^2 E_y = 0$$

$$E_y = 0 \quad \text{at} \quad x = 0 \quad \text{and} \quad x = a - d$$

where d is assumed to be small. Compare the result with the exact value $k^2 = \pi^2 / (a - d)^2$.

- 3.2 Prove that (3.96) is a variational expression for the eigenvalue problem defined by (3.93) and (3.94).

- 3.3 Using the variational expression (3.96), calculate the eigenvalue of the TE_{10} -like mode in a waveguide whose cross section is a rectangle with the four corners removed as shown in Fig. 3.18.

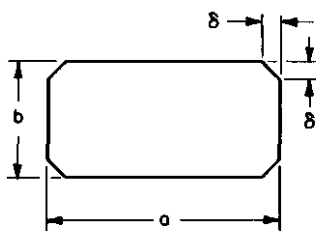


Fig. 3.18. Deformed cross section of a rectangular waveguide.

- 3.4 Calculate the electric and magnetic fields of the dominant mode in a semicircular waveguide.
- 3.5 The rectangular TE_{11} and TM_{11} modes are degenerate to each other. Following the discussion in Section 3.8, calculate how the removal of the degeneracy takes place when the wall losses are taken into account.
- 3.6 Suppose the inner surface of a circular waveguide with radius a is coated by dielectric material ($\epsilon = \epsilon_0 \epsilon_s$) of thickness δ ($\delta \ll a$). Discuss the effect of the coating on the TE_{01} and TM_{11} modes by calculating how the degeneracy is removed.
- 3.7 Suppose two rectangular waveguides are placed side by side and two holes are drilled through the common wall a quarter wavelengths apart as shown in Fig. 3.19. Part of the power coming in from port I will leak out from port IV, but none will emerge

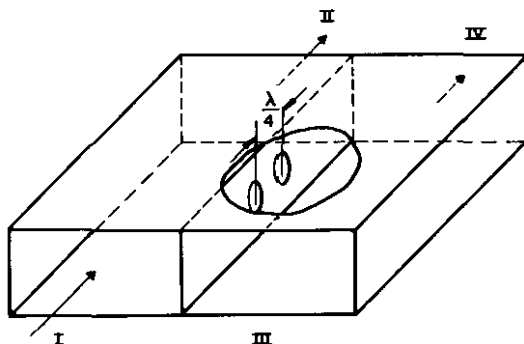


Fig. 3.19. Directional coupler.

from port III since the waves coming through the two holes cancel each other due to the half wavelength path difference. Similarly, if a wave is incident on port II, some power leaks out from port III, but none emerges from port IV. Because of this directional property, the device is called the directional coupler. If three holes are drilled, each a quarter wavelength apart, and the amount of leakage through each hole has the voltage ratio 1:2:1, then the directional property is obtained over a wider frequency bandwidth than in the case with only two holes. Using a complex vector diagram with three vectors, each representing the magnitude and phase of the leakage through each hole, explain why the bandwidth becomes wider.

- 3.8 Cross one rectangular waveguide on top of another and drill two cross-shaped holes through the common wall as shown in Fig. 3.20. Then, part of the power coming in

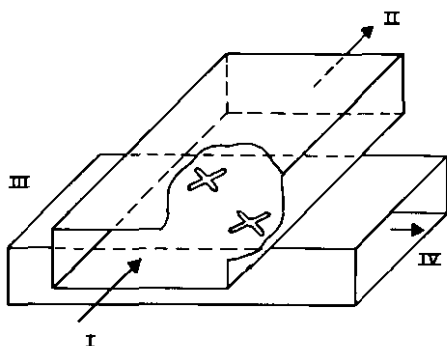


Fig. 3.20. Cross-guide directional coupler.

from port I leaks out from port IV but none from port III, giving the directional coupler action. Explain the principle of operation qualitatively, considering the field configuration of the rectangular TE_{10} mode. This is inherently a broadband device since the cancellation does not utilize the path length difference.

- 3.9 Prove that γ^2 is real for the waveguides with inhomogeneous media when $\sigma = 0$. The E_{tm} 's can be chosen to be real, and (3.193) is equivalent to

$$\int \mathbf{k} \cdot \mathbf{E}_{tm} \times \mathbf{H}_{tn}^* dS = 0$$

for the propagating modes m and n , which indicates that each wave carries its own independent power.

- 3.10 Prove that a variational expression for the eigenvalue problem $(d^2 E_y/dx^2) + k^2 E_y = 0$ under the boundary conditions $E_y(0) = E_y(a)$ and $\{dE_y(0)/dx\} = \{dE_y(a)/dx\}$ is given by

$$k^2(E_y) =$$

$$\frac{\int (dE_y/dx)^2 dx - [E_y dE_y/dx]_0^a + E_y(0) \{dE_y(a)/dx\} - E_y(a) \{dE_y(0)/dx\}}{\int E_y^2 dx}$$

CHAPTER 4

RESONANT CAVITIES

Resonant cavities are devices constructed so that one can utilize resonant phenomena of electromagnetic fields in a space enclosed by good conducting walls. They are useful as circuit elements, particularly as microwave counterparts of LC resonant circuits in low frequency ranges. Furthermore, their analysis offers an opportunity to demonstrate a powerful eigenfunction approach. For these reasons, this whole chapter is devoted to the theory of resonant cavities. An equivalent circuit of the cavities is obtained and a method to experimentally determine the circuit parameters is studied in detail. In addition, a discussion of cavities with inhomogeneous media is briefly presented. In contrast to waveguides with inhomogeneous media, the completeness of the eigenfunctions can be shown without difficulty.

4.1 Introduction

Conductor walls forming a resonant cavity generally have finite conductivities and, hence, introduce losses of electromagnetic energy. Further, in order to utilize the resonances in the cavity, there must be at least one opening through which the inside and outside of the cavity are connected. However, as an idealized model, let us first consider a space completely enclosed by a perfect conductor and study how electromagnetic fields behave in it. Maxwell's equations are given by

$$\nabla \times \mathbf{H} = j\omega\epsilon\mathbf{E} \quad (4.1)$$

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H} \quad (4.2)$$

Substituting $\nabla \times (4.2)$ into (4.1), we have an equation in $\mathbf{E}$ alone;

$$\nabla \times \nabla \times \mathbf{E} - \omega^2 \epsilon \mu \mathbf{E} = 0 \quad (\text{in } V) \quad (4.3)$$

where an assumption is made that ϵ and μ are constant in the volume V of the cavity. The boundary condition for $\mathbf{E}$ is given by

$$\mathbf{n} \times \mathbf{E} = 0 \quad (\text{on } S) \quad (4.4)$$

where $\mathbf{n}$ is the unit vector directed outwards normal to the wall surface S . Equations (4.3) and (4.4) constitute another eigenvalue problem whose solutions can exist only when $\omega^2 \epsilon \mu = k^2$ takes certain discrete values. For instance, for a rectangular enclosure with sides a , b and c , the eigenvalues are given by

$$k_a^2 = \omega_a^2 \epsilon \mu = (n\pi/a)^2 + (m\pi/b)^2 + (l\pi/c)^2 \quad (4.5)$$

where n , m , and l are arbitrary integers including zero, provided two of them do not become zero simultaneously. This type of cavity can be considered to be a rectangular waveguide which is an integer multiple of a half wavelength long and short circuited at both ends. There are two different field configurations for each set of n , m , and l , one corresponding to a TE mode, and the other to a TM mode. Accordingly, the field configurations are designated as TE_{nml} and TM_{nml} modes. There is one exception; no TE mode exists for $l = 0$.

For a cylindrical cavity of radius a and length L , corresponding to TE modes in a cylindrical waveguide, there are TE_{nm} modes ($l \neq 0$) whose eigenvalues are given by

$$k_a^2 = \omega_a^2 \epsilon \mu = (U'_{nm}/a)^2 + (l\pi/L)^2 \quad (4.6)$$

Similarly, corresponding to TM modes of the waveguide, there are TM_{nm} modes with the eigenvalues given by

$$k_a^2 = \omega_a^2 \epsilon \mu = (U_{nm}/a)^2 + (l\pi/L)^2 \quad (4.7)$$

where l is allowed to be equal to zero. The resonant frequencies for $l = 0$ are exactly the cutoff frequencies of the corresponding modes in the waveguide.

As illustrated above, if a finite electromagnetic field exists in a space completely enclosed by a perfect conductor, the field should have at least one of the discrete frequencies determined by $\omega_a = k_a(\epsilon\mu)^{-1/2}$ where the k_a 's are the eigenvalues of (4.3) and (4.4).

In order to excite, detect, or utilize the electromagnetic field in the cavity,

let us next consider two small openings as shown in Fig. 4.1, which connect the inside of the cavity to waveguide ports I and II. Suppose a wave with frequency ω is incident on port I, and ω is equal to one of the discrete frequencies mentioned above, then the electromagnetic field will be strongly excited in the cavity and a substantial amount of power will leak out to port II. On the other hand, if ω is different from any one of the ω_a 's, port I and port II will be isolated by the cavity since it cannot contain a field with

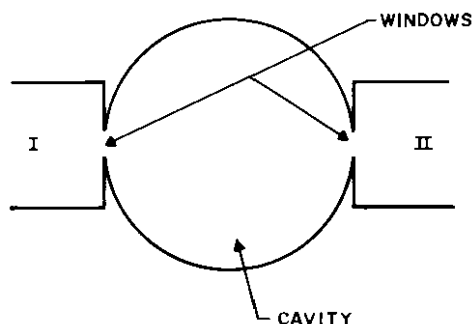


Fig. 4.1. Resonant cavity with coupling windows.

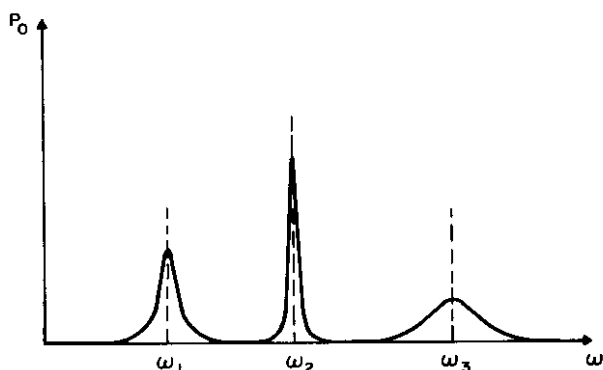


Fig. 4.2. Cavity output power versus ω .

ω different from the ω_a 's. This, however, is true only for the idealized cavity with no opening. Because of the opening, if the difference between ω and ω_a is small, a weak electromagnetic field will be excited in the cavity and some power will leak out of port II. As a result, when ω is varied keeping the amplitude of the wave incident on port I constant, the power out of port II may look like Fig. 4.2. At ω_1 , ω_2 , ω_3 , and so on, determined by the geometry of the cavity, the output power becomes large and as ω moves away from

these discrete frequencies, it drops rapidly. Each of the bell shaped curves is called a resonant curve. The heights of the peaks are different from one another because different modes are excited differently. When the cavity is close to the idealized case, the output power drops quickly as ω changes from each ω_a , and the width of the resonant curve becomes narrow. On the other hand, if the resonant curve is broad, it shows that the cavity behaves quite differently from the idealized case for the particular mode. Thus, the width of the resonant curve gives some indication about the quality of the cavity and how close it is to being ideal. The resonant curve is sometimes called the Q -curve after the first letter of the word "quality." The center frequency of the curve is equal to ω_a (for a more precise discussion, see Sections 4.2 and 4.3), a property which can be used to measure frequencies when we change the length of the cavity to make the output power maximum. If either the approximate frequency is known beforehand or the cavity is constructed so that only one mode can be excited easily, the resonant mode can be identified, and the frequency ω_a can be calculated from the dimensions of the cavity. Cylindrical cavities are widely used for this purpose since it is relatively easy to build them precisely using a lathe. The modes usually employed for frequency measurements are circular TE_{01l} or TE_{11l} where l is a small integer such as 1, 2, or 3.

4.2 Expansions of Electromagnetic Fields

A complete set of orthonormal functions cannot be derived from the eigenvalue problem defined by (4.3) and (4.4); nevertheless, the two-dimensional problem studied in the previous chapter suggests that the solutions of the eigenvalue problem

$$\nabla \times \nabla \times \mathbf{E} - \nabla \nabla \cdot \mathbf{E} - k^2 \mathbf{E} = 0 \quad (\text{in } V) \quad (4.8)$$

$$\mathbf{n} \times \mathbf{E} = 0, \quad \nabla \cdot \mathbf{E} = 0 \quad (\text{on } S) \quad (4.9)$$

may give a desired set; in fact, this eigenvalue problem can be discussed in exactly the same manner as before. Thus, the eigenvalues are real and non-negative (positive or zero), and the real eigenfunctions can be chosen to form a complete set. The variational expression for k^2 is given by

$$k^2(\mathbf{E}) = \frac{\int \{(\nabla \times \mathbf{E})^2 + (\nabla \cdot \mathbf{E})^2\} dv - 2 \int \mathbf{n} \times \mathbf{E} \cdot \nabla \times \mathbf{E} dS}{\int \mathbf{E}^2 dv} \quad (4.10)$$

where the volume integral and surface integral are over V and S , respectively. The orthogonality and normalization conditions are expressed by

$$\begin{aligned}\int \mathbf{E}_n \cdot \mathbf{E}_m dv &= 0 & (n \neq m) \\ \int \mathbf{E}_n \cdot \mathbf{E}_m dv &= 1 & (n = m)\end{aligned}$$

Furthermore, each function in the complete set can be chosen so as to belong to one of the following three groups.

$$\begin{aligned}\text{I. } \nabla \times \mathbf{E}_n &= 0, & \nabla \cdot \mathbf{E}_n &= 0 \\ \text{II. } \nabla \times \mathbf{E}_n &\neq 0, & \nabla \cdot \mathbf{E}_n &= 0 \\ \text{III. } \nabla \times \mathbf{E}_n &= 0, & \nabla \cdot \mathbf{E}_n &\neq 0\end{aligned}$$

In terms of the eigenfunctions thus obtained, an arbitrary vector function $\mathbf{F}$ can be expanded as follows, provided it is piecewise-continuous (i.e., continuous except on a finite number of surfaces each with a finite area) and square-integrable over V :

$$\mathbf{F} = \sum_{n=1}^{\infty} A_n \mathbf{E}_n \quad (4.11)$$

where

$$A_n = \int \mathbf{F} \cdot \mathbf{E}_n dv \quad (4.12)$$

The real meaning of the equality in (4.11) is given by

$$\lim_{N \rightarrow \infty} \int \left(\mathbf{F} - \sum_{n=1}^{N-1} A_n \mathbf{E}_n \right)^2 dv = 0 \quad (4.13)$$

When $\mathbf{F}$ is a complex vector function, the real and imaginary parts can be expanded separately and then added to obtain the formulas identical to (4.11) and (4.12). No proofs will be given here for the above assertions since they are obvious from the previous discussion for the two-dimensional case. A brief discussion concerning the proof of the relation $\lim_{n \rightarrow \infty} k_n^2 = \infty$ will, however, be given in Appendix I.

Following a procedure similar to the eigenvalue problem for $\mathbf{H}_t$ in the previous chapter, we have

$$\nabla \times \nabla \times \mathbf{H} - \nabla \nabla \cdot \mathbf{H} - k^2 \mathbf{H} = 0 \quad (\text{in } V) \quad (4.14)$$

$$\mathbf{n} \cdot \mathbf{H} = 0, \quad \mathbf{n} \times \nabla \times \mathbf{H} = 0 \quad (\text{on } S) \quad (4.15)$$

A variational expression for k^2 is given by

$$k^2(\mathbf{H}) = \frac{\int \{(\nabla \times \mathbf{H})^2 + (\nabla \cdot \mathbf{H})^2\} dv - 2 \int (\mathbf{n} \cdot \mathbf{H})(\nabla \cdot \mathbf{H}) dS}{\int \mathbf{H}^2 dv} \quad (4.16)$$

The orthogonality and normalization conditions are

$$\int \mathbf{H}_n \cdot \mathbf{H}_m dv = 0 \quad (n \neq m), \quad \int \mathbf{H}_n \cdot \mathbf{H}_m dv = 1 \quad (n = m)$$

Furthermore, we shall assume that the eigenfunctions are real and belong to one of the following three groups:

- I. $\nabla \times \mathbf{H}_m = 0, \quad \nabla \cdot \mathbf{H}_m = 0$
- II. $\nabla \times \mathbf{H}_m \neq 0, \quad \nabla \cdot \mathbf{H}_m = 0$
- III. $\nabla \times \mathbf{H}_m = 0, \quad \nabla \cdot \mathbf{H}_m \neq 0$

A piecewise-continuous and square-integrable, but otherwise arbitrary, vector function $\mathbf{F}$, defined in V , can be expanded in terms of these functions:

$$\mathbf{F} = \sum_{m=1}^{\infty} B_m \mathbf{H}_m \quad (4.17)$$

where

$$B_m = \int \mathbf{F} \cdot \mathbf{H}_m dv \quad (4.18)$$

We have obtained two complete sets of orthonormal functions, the $\mathbf{E}_n$'s and the $\mathbf{H}_m$'s, but those belonging to group II are closely related to each other. To show this, we first note that their eigenvalues are not equal to zero, as can be seen from the variational expressions. Using $\mathbf{E}_a$, which is one of the $\mathbf{E}_n$'s belonging to group II, we can define a function $\mathbf{H}_a$ through

$$\nabla \times \mathbf{E}_a = k_a \mathbf{H}_a \quad (4.19)$$

where k_a is the nonzero eigenvalue of $\mathbf{E}_a$. Therefore, $\mathbf{H}_a$ belongs to group II and can be adapted as one of the $\mathbf{H}_n$'s to form a complete set for the following reasons. Since $\nabla \cdot \mathbf{H}_a = 0$, $\nabla \cdot \mathbf{E}_a = 0$, and

$$\nabla \times \nabla \times \mathbf{H}_a - k_a^2 \mathbf{H}_a = k_a^{-1} \nabla \times (\nabla \times \nabla \times \mathbf{E}_a - k_a^2 \mathbf{E}_a) = 0 \quad (\text{in } V)$$

the function $\mathbf{H}_a$ satisfies (4.14). On S , it satisfies

$$\mathbf{n} \times \nabla \times \mathbf{H}_a = k_a^{-1} \mathbf{n} \times \nabla \times \nabla \times \mathbf{E}_a = k_a^{-1} \mathbf{n} \times k_a^2 \mathbf{E}_a = 0 \quad (\text{on } S)$$

which is one of the boundary conditions in (4.15). For the other boundary condition, let us consider the surface integral of $\mathbf{n} \cdot \mathbf{H}_a$ over ΔS , an arbitrary part of S . It is given by

$$\int_{\Delta S} \mathbf{n} \cdot \mathbf{H}_a dS = k_a^{-1} \int_{\Delta S} \mathbf{n} \cdot \nabla \times \mathbf{E}_a dS = k_a^{-1} \int_{\Delta L} \mathbf{E}_a \cdot d\mathbf{l}$$

where ΔL is a closed contour around ΔS . Since the component of $\mathbf{E}_a$ tangential to S is equal to zero, the right-hand side vanishes regardless of the size and shape of ΔS . This means that

$$\mathbf{n} \cdot \mathbf{H}_a = 0 \quad (\text{on } S)$$

Thus, $\mathbf{H}_a$ satisfies (4.14) and (4.15). Furthermore, for $\mathbf{H}_a$ and $\mathbf{H}_b$ corresponding to $\mathbf{E}_a$ and $\mathbf{E}_b$, respectively, we have

$$\begin{aligned} \int \mathbf{H}_a \cdot \mathbf{H}_b dv &= (k_a k_b)^{-1} \int \nabla \times \mathbf{E}_a \cdot \nabla \times \mathbf{E}_b dv \\ &= (k_a k_b)^{-1} \int \mathbf{n} \times \mathbf{E}_a \cdot \nabla \times \mathbf{E}_b dS + (k_b/k_a) \int \mathbf{E}_a \cdot \mathbf{E}_b dv \end{aligned}$$

The first term on the right-hand side is equal to zero since $\mathbf{n} \times \mathbf{E}_a = 0$. If $a \neq b$, the second term also vanishes due to the orthogonality relation between $\mathbf{E}_a$ and $\mathbf{E}_b$. If $a = b$, the integral becomes unity. These two conclusions mean that the $\mathbf{H}_a$'s satisfy the orthonormal conditions when the $\mathbf{E}_a$'s do. We have seen that an eigenfunction $\mathbf{H}_a$ belonging to group II can be derived from each eigenfunction $\mathbf{E}_a$ in group II. Conversely, if $\mathbf{H}_a$ belongs to group II, $\mathbf{E}_a$ can be defined through

$$\nabla \times \mathbf{H}_a = k_a \mathbf{E}_a \quad (4.20)$$

As a result, every eigenfunction belonging to group II can be made to have a counterpart in the other set without loss of completeness. Let us assume that this has been done in the remainder of the discussion.

In order to distinguish eigenfunctions in groups I and III from those in group II, Greek subscripts will be used, i.e.;

$$\nabla \times \mathbf{E}_\nu = 0 \quad (4.21)$$

$$\nabla \times \mathbf{H}_\lambda = 0 \quad (4.22)$$

We are now in a position to solve Maxwell's equations using a strategy similar to the one employed in Section 3.5. First, we shall expand all the quantities appearing in Maxwell's equations in terms of appropriate sets of

functions obtained above. We shall then determine the coefficients by substituting the results into Maxwell's equations, thereby obtaining the electric and magnetic fields in expanded forms.

If there were no openings or losses in the walls, the $\mathbf{E}_a$'s and the $\mathbf{H}_a$'s would suffice to define the electric and magnetic fields. However, because of these imperfections, all the functions, the $\mathbf{E}_a$'s, $\mathbf{E}_v$'s, $\mathbf{H}_a$'s, and $\mathbf{H}_v$'s, become necessary for our discussion. For simplicity, let us first consider a cavity with only one opening connected to a waveguide. In order to facilitate the connection of the fields in the cavity with those in the waveguide, let us extend the cavity region V into the waveguide, as shown in Fig. 4.3 where

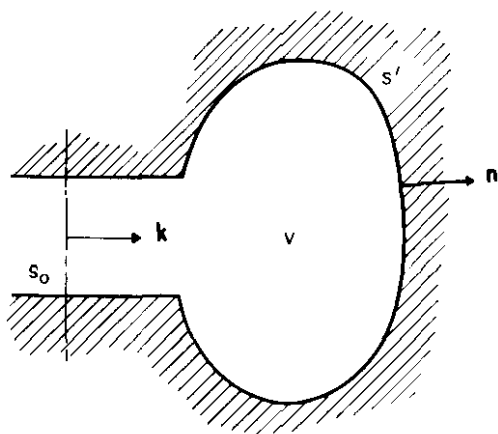


Fig. 4.3. Cavity volume V is extended to reference plane S_0 .

S_0 is a cross section of the waveguide defining the extent of V . Let S' be the cavity wall surface so that $S' + S_0$ forms a closed surface S completely enclosing V .

Maxwell's equations, which we are going to solve, are given by

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H} \quad (4.23)$$

$$\nabla \times \mathbf{H} = (\sigma + j\omega\epsilon)\mathbf{E} \quad (4.24)$$

On the right-hand side of (4.24), σ has been introduced for several reasons. Although we shall assume that $\sigma \ll j\omega\epsilon$, if σ is finite, the effect of σ becomes important at and near the resonant frequencies and hence cannot be neglected. Furthermore, use will be made of the fact that $k_a^2 = \omega_a^2\epsilon\mu$ is real and, hence, the replacement of ϵ by $\epsilon + \sigma/j\omega$ in the final result will not give a correct answer as it did in the case of waveguides.

The expansions of $\mathbf{E}$ and $\mathbf{H}$ are given by

$$\mathbf{E} = \sum_a \mathbf{E}_a \int \mathbf{E} \cdot \mathbf{E}_a dv + \sum_v \mathbf{E}_v \int \mathbf{E} \cdot \mathbf{E}_v dv \quad (4.25)$$

$$\mathbf{H} = \sum_a \mathbf{H}_a \int \mathbf{H} \cdot \mathbf{H}_a dv + \sum_\lambda \mathbf{H}_\lambda \int \mathbf{H} \cdot \mathbf{H}_\lambda dv \quad (4.26)$$

Next, considering that $\nabla \times \mathbf{E}$ is an $\mathbf{H}$ -like function, let us expand it in the form

$$\nabla \times \mathbf{E} = \sum_a \mathbf{H}_a \int \nabla \times \mathbf{E} \cdot \mathbf{H}_a dv + \sum_\lambda \mathbf{H}_\lambda \int \nabla \times \mathbf{E} \cdot \mathbf{H}_\lambda dv \quad (4.27)$$

The first expansion coefficients on the right-hand side can be rewritten as follows:

$$\begin{aligned} \int \nabla \times \mathbf{E} \cdot \mathbf{H}_a dv &= \int (\mathbf{E} \cdot \nabla \times \mathbf{H}_a + \nabla \cdot \mathbf{E} \times \mathbf{H}_a) dv \\ &= k_a \int \mathbf{E} \cdot \mathbf{E}_a dv + \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_a dS \end{aligned}$$

Similarly, since $\nabla \times \mathbf{H}_\lambda = 0$, we have

$$\int \nabla \times \mathbf{E} \cdot \mathbf{H}_\lambda dv = \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_\lambda dS$$

Substituting these expressions into (4.27), the expansion of $\nabla \times \mathbf{E}$ becomes

$$\nabla \times \mathbf{E} = \sum_a \mathbf{H}_a \left(k_a \int \mathbf{E} \cdot \mathbf{E}_a dv + \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_a dS \right) + \sum_\lambda \mathbf{H}_\lambda \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_\lambda dS \quad (4.28)$$

If we took $\nabla \times$ (4.25) and calculated the term-by-term differentiation of the right-hand side, the surface integrals in (4.28) would not appear, thus leading to an incorrect result. This is due to the fact that the term-by-term differentiation of an infinite series is not always allowable, especially when the boundary conditions of the functions in the series expansion are different from those of the function being expanded.

Finally, let us expand $\nabla \times \mathbf{H}$ in terms of the $\mathbf{E}_a$'s and $\mathbf{E}_v$'s. A little manipulation similar to that leading to (4.28) gives

$$\nabla \times \mathbf{H} = \sum_a \mathbf{E}_a k_a \int \mathbf{H} \cdot \mathbf{H}_a dv \quad (4.29)$$

The surface integrals corresponding to the ones in (4.28) do not appear because $\mathbf{n} \times \mathbf{E}_a$ and $\mathbf{n} \times \mathbf{E}_v$ are both equal to zero on S .

Since all the necessary quantities have been obtained in their expanded forms, let us substitute them into (4.23) and (4.24). The results are given by

$$\begin{aligned} \sum_a \mathbf{H}_a \left(k_a \int \mathbf{E} \cdot \mathbf{E}_a dv + \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_a dS \right) + \sum_\lambda \mathbf{H}_\lambda \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_\lambda dS \\ = -j\omega\mu \left(\sum_a \mathbf{H}_a \int \mathbf{H} \cdot \mathbf{H}_a dv + \sum_\lambda \mathbf{H}_\lambda \int \mathbf{H} \cdot \mathbf{H}_\lambda dv \right) \end{aligned} \quad (4.30)$$

$$\begin{aligned} \sum_a \mathbf{E}_a k_a \int \mathbf{H} \cdot \mathbf{H}_a dv \\ = (\sigma + j\omega\epsilon) \left(\sum_a \mathbf{E}_a \int \mathbf{E} \cdot \mathbf{E}_a dv + \sum_v \mathbf{E}_v \int \mathbf{E} \cdot \mathbf{E}_v dv \right) \end{aligned} \quad (4.31)$$

Equating the coefficients of each eigenfunction on both sides of these equations we obtain

$$k_a \int \mathbf{E} \cdot \mathbf{E}_a dv + \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_a dS = -j\omega\mu \int \mathbf{H} \cdot \mathbf{H}_a dv \quad (4.32)$$

$$\int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_\lambda dS = -j\omega\mu \int \mathbf{H} \cdot \mathbf{H}_\lambda dv \quad (4.33)$$

$$k_a \int \mathbf{H} \cdot \mathbf{H}_a dv = (\sigma + j\omega\epsilon) \int \mathbf{E} \cdot \mathbf{E}_a dv \quad (4.34)$$

$$\int \mathbf{E} \cdot \mathbf{E}_v dv = 0 \quad (4.35)$$

If we eliminate one of the volume integrals from (4.32) and (4.34), say $\int \mathbf{E} \cdot \mathbf{E}_a dv$, we obtain

$$\{j\omega\mu + k_a^2/(\sigma + j\omega\epsilon)\} \int \mathbf{H} \cdot \mathbf{H}_a dv = - \int \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_a dS \quad (4.36)$$

Now suppose $\mathbf{n} \times \mathbf{E}$ is given on S , then $\int \mathbf{H} \cdot \mathbf{H}_a dv$ can be calculated from (4.36). Substituting this into (4.34), $\int \mathbf{E} \cdot \mathbf{E}_a dv$ will be determined. Furthermore, $\int \mathbf{H} \cdot \mathbf{H}_\lambda dv$ can be calculated from (4.33) while $\int \mathbf{E} \cdot \mathbf{E}_v dv$ is equal to zero from (4.35). Substituting all these volume integrals into (4.25) and (4.26), the electric and magnetic fields are obtained in their expanded forms and Maxwell's equations are solved assuming $\mathbf{n} \times \mathbf{E}$ is known. To complete the solution our next task is to find $\mathbf{n} \times \mathbf{E}$ on $S = S_0 + S'$. First consider $\mathbf{n} \times \mathbf{E}$ on S_0 . Since S_0 is a cross section of the waveguide, the component $\mathbf{E}_\parallel$ of $\mathbf{E}$ tangential to S_0 can be expanded in terms of the eigenfunctions in the

waveguide:

$$\mathbf{E}_{\parallel} = \sum_n \mathbf{E}_{in} V_n \quad (4.37)$$

where V_n is the expansion coefficient which can be interpreted as the voltage associated with the n th mode in the waveguide as discussed in Section 3.5.

Next, to calculate the surface integral on the right-hand side of (4.36), let us expand $-\mathbf{k} \times \mathbf{H}_a$ on S_0 in terms of the $\mathbf{E}_{in}$'s:

$$-\mathbf{k} \times \mathbf{H}_a = \sum_n \mathbf{E}_{in} I_{an}$$

where I_{an} is the expansion coefficient. Multiplying both sides by $\mathbf{k} \times$ and using the relation

$$\mathbf{k} \times \mathbf{k} \times \mathbf{H}_a = \mathbf{k}(\mathbf{k} \cdot \mathbf{H}_a) - \mathbf{H}_a(\mathbf{k} \cdot \mathbf{k}) = -\mathbf{H}_a$$

we have

$$\mathbf{H}_a = \sum_n \mathbf{k} \times \mathbf{E}_{in} I_{an} \quad (4.38)$$

Noting that the direction of $\mathbf{k}$ is opposite to that of $\mathbf{n}$, we obtain from (4.37) and (4.38)

$$\begin{aligned} - \int_{S_0} \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_a dS &= \int_{S_0} \mathbf{k} \times \mathbf{E} \cdot \mathbf{H}_a dS \\ &= \int_{S_0} \sum_n \mathbf{k} \times \mathbf{E}_{in} V_n \cdot \sum_n \mathbf{k} \times \mathbf{E}_{in} I_{an} dS \\ &= \sum_n V_n I_{an} \end{aligned} \quad (4.39)$$

where a formula similar to (3.85) is used. In much the same way, $\mathbf{H}_\lambda$ can be expressed in the form

$$\mathbf{H}_\lambda = \sum_n \mathbf{k} \times \mathbf{E}_{in} I_{\lambda n} \quad (4.40)$$

and the contribution from S_0 to the surface integral on the left-hand side of (4.33) becomes

$$\int_{S_0} \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_\lambda dS = - \sum_n V_n I_{\lambda n} \quad (4.41)$$

Having obtained the integrals over S_0 , let us now turn our attention to the wall surface S' . For simplicity, we assume that ω is close to $\omega_p = k_p(\epsilon\mu)^{-1/2}$ where k_p is one of the k_a 's and that only the p th mode is strongly excited. All the other terms in the field expansions will be considered small compared to this mode. Then $\mathbf{H}$ can be replaced by $\mathbf{H}_p \int \mathbf{H} \cdot \mathbf{H}_p dv$ in order

to calculate the effect of the wall losses. As we discussed in Section 3.8, the boundary condition on S' is given by

$$\mathbf{n} \times \mathbf{E} = Z_w \mathbf{H}$$

and, hence, we have

$$\mathbf{n} \times \mathbf{E} = Z_w \mathbf{H}_p \int \mathbf{H} \cdot \mathbf{H}_p dv \quad (4.42)$$

under the assumption mentioned above. From (4.42), the surface integral of $\mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_p$ over S' becomes

$$\int_{S'} \mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_p dS = (1+j) \omega_p \mu Q_{sp}^{-1} \int \mathbf{H} \cdot \mathbf{H}_p dv \quad (4.43)$$

where

$$Q_{sp}^{-1} = (1+j)^{-1} (\omega_p \mu)^{-1} \int_{S'} Z_w \mathbf{H}_p^2 dS \quad (4.44)$$

Substituting (4.39) and (4.43) into (4.36), with $a = p$ and assuming $j\omega\epsilon \gg \sigma$, we obtain

$$\{k_p^2 - \omega^2 \epsilon \mu + j\omega \mu \sigma + (1+j)j\omega\epsilon \omega_p \mu Q_{sp}^{-1}\} \int \mathbf{H} \cdot \mathbf{H}_p dv = j\omega\epsilon \sum_n V_n I_{pn} \quad (4.45)$$

Dividing both sides by $j\omega\omega_p \epsilon \mu$ and rearranging the terms, when ω is close to ω_p the above equation reduces to

$$\int \mathbf{H} \cdot \mathbf{H}_p dv = \frac{1}{\omega_p \mu j \{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \sum_n V_n I_{pn} \quad (4.46)$$

where

$$\omega_p' = \omega_p (1 - \frac{1}{2} Q_{sp}^{-1}), \quad (1/Q_p') = (1/Q_p) + (1/Q_{sp}), \quad (1/Q_p) = (\sigma/\omega_p \epsilon) \quad (4.47)$$

Since the first and second terms in the brackets on the left-hand side of (4.45) almost cancel each other, the contribution from (4.43) to the left-hand side of (4.45) cannot be neglected in obtaining (4.46). On the other hand, since Z_w is small, the contribution from the surface integral of $\mathbf{n} \times \mathbf{E} \cdot \mathbf{H}_a$ ($a \neq p$) over S' is negligible compared to the other terms in (4.36). Thus, we have

$$\int \mathbf{H} \cdot \mathbf{H}_a dv = \frac{1}{\omega_a \mu j \{(\omega/\omega_a) - (\omega_a/\omega)\} + (1/Q_a)} \sum_n V_n I_{an} \quad (4.48)$$

where a is not equal to p . Similarly, we obtain from (4.33) and (4.41)

$$\int \mathbf{H} \cdot \mathbf{H}_\lambda dv = \frac{\sum_n V_n I_{\lambda n}}{j\omega\mu} \quad (4.49)$$

Substituting (4.46), (4.48), and (4.49) into (4.26), the magnetic field in the cavity is obtained:

$$\begin{aligned} \mathbf{H} = & \mathbf{H}_p \frac{1}{\omega_p \mu j} \frac{\sum_n V_n I_{pn}}{\{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \\ & + \sum_{a \neq p} \mathbf{H}_a \frac{1}{\omega_a \mu j} \frac{\sum_n V_n I_{an}}{\{(\omega/\omega_a) - (\omega_a/\omega)\} + (1/Q_a)} + \sum_\lambda \mathbf{H}_\lambda \frac{\sum_n V_n I_{\lambda n}}{j\omega\mu} \end{aligned} \quad (4.50)$$

Similarly, $\mathbf{E}$ is given by

$$\begin{aligned} \mathbf{E} = & -j \left(\frac{\mu}{\epsilon} \right)^{1/2} \left[\mathbf{E}_p \frac{1}{\omega_p \mu j} \frac{\sum_n V_n I_{pn}}{\{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \right. \\ & \left. + \sum_{a \neq p} \mathbf{E}_a \frac{1}{\omega_a \mu j} \frac{\sum_n V_n I_{an}}{\{(\omega/\omega_a) - (\omega_a/\omega)\} + (1/Q_a)} \right] \end{aligned} \quad (4.51)$$

This completes the solution of Maxwell's equations.

Let us next calculate the input admittance of the cavity, looking in from S_0 . To do so, we assume that S_0 is far from any waveguide discontinuities and that only one propagating waveguide mode exists there; i.e., all the other modes are negligibly small. From (4.38), (4.40), and (4.50), the tangential component of $\mathbf{H}$ on S_0 is given by

$$\begin{aligned} \mathbf{H}_\parallel = & \sum_n \sum_m \mathbf{k} \times \mathbf{E}_{tm} I_{pm} \frac{1}{\omega_p \mu j} \frac{V_n I_{pn}}{\{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} + \sum_n \sum_m \mathbf{k} \times \mathbf{E}_{tm} \\ & \times \left[\sum_{a \neq p} I_{am} \frac{1}{\omega_a \mu j} \frac{V_n I_{an}}{\{(\omega/\omega_a) - (\omega_a/\omega)\} + (1/Q_a)} + \sum_\lambda I_{\lambda m} \frac{V_n I_{\lambda n}}{j\omega\mu} \right] \end{aligned} \quad (4.52)$$

If subscript 1 indicates the propagating mode in (4.52), the terms other than those corresponding to $m = n = 1$ must all be negligible, by hypothesis. On the other hand, the same magnetic field can be written in the form

$$\mathbf{H}_\parallel = \mathbf{k} \times \mathbf{E}_{t1} I_1 \quad (4.53)$$

where I_1 is the current associated with the propagating mode. Therefore,

I_1 is given by

$$I_1 = \frac{1}{\omega_p \mu j \{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \frac{V_1 I_{p1}^2}{\omega_p \mu j \{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} + \sum_{a \neq p} \frac{1}{\omega_a \mu j \{(\omega/\omega_a) - (\omega_a/\omega)\} + (1/Q_a)} \frac{V_1 I_{a1}^2}{j\omega \mu} + \sum_{\lambda} \frac{V_1 I_{\lambda 1}^2}{j\omega \mu}$$

Since V_1 is the voltage associated with the propagating mode in the waveguide, the input admittance of the cavity becomes

$$Y = \frac{I_1}{V_1} = \frac{1}{\omega_p \mu j \{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \frac{I_{p1}^2}{V_1} + Y_p \quad (4.54)$$

where

$$Y_p = \sum_{a \neq p} \frac{1}{\omega_a \mu j \{(\omega/\omega_a) - (\omega_a/\omega)\} + (1/Q_a)} \frac{I_{a1}^2}{V_1} + \sum_{\lambda} \frac{I_{\lambda 1}^2}{j\omega \mu} \quad (4.55)$$

The term Y_p is a slowly varying function of ω in the vicinity of $\omega = \omega_p$ in contrast to the first term on the right-hand side of (4.54).

Corresponding to (4.54), for resonant cavities with two openings, such as that illustrated in Fig. 4.1, we obtain

$$I_1 = Y_{11}V_1 + Y_{12}V_2, \quad I_2 = Y_{12}V_1 + Y_{22}V_2 \quad (4.56)$$

where

$$Y_{ij} = \frac{1}{\omega_p \mu j \{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \frac{I_{pi} I_{pj}}{V_1} + \sum_{a \neq p} \frac{1}{\omega_a \mu j \{(\omega/\omega_a) - (\omega_a/\omega)\} + (1/Q_a)} \frac{I_{ai} I_{aj}}{j\omega \mu} + \sum_{\lambda} \frac{I_{\lambda i} I_{\lambda j}}{j\omega \mu} \quad (i, j = 1, 2) \quad (4.57)$$

Equations (4.54) and (4.56) will be studied in detail in the next section.

The electric and magnetic fields of the p th mode in the cavity, $[\mathbf{E}]_p$ and $[\mathbf{H}]_p$, are given by

$$[\mathbf{E}]_p = -j \left(\frac{\mu}{\epsilon} \right)^{1/2} \frac{1}{\omega_p \mu j \{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \frac{\sum_n V_n I_{pn}}{V_1} \mathbf{E}_p$$

$$[\mathbf{H}]_p = \frac{1}{\omega_p \mu j \{(\omega/\omega_p') - (\omega_p'/\omega)\} + (1/Q_p')} \frac{\sum_n V_n I_{pn}}{V_1} \mathbf{H}_p$$

Therefore, the relation

$$\frac{1}{2} \int \epsilon [\mathbf{E}]_p \cdot [\mathbf{E}]_p^* dv = \frac{1}{2} \int \mu [\mathbf{H}]_p \cdot [\mathbf{H}]_p^* dv \quad (4.58)$$

holds at $\omega = \omega_p$. This means that the average electric and magnetic stored energies of a resonant mode are equal at the resonant frequency. Since $\mathbf{E}_p$ and $\mathbf{H}_p$ are real, $[\mathbf{E}]_p$ and $[\mathbf{H}]_p$ are 90° out of phase. When the electric stored energy becomes maximum, the magnetic stored energy becomes zero and vice versa. However, the sum of these stored energies always remains constant at the resonant frequency. This sum is twice as large as the average electric or magnetic stored energy.

As defined in (4.47), Q_p' has the same physical meaning as the ordinary Q of a resonant circuit in low frequency ranges. This can be shown by re-writing $1/Q_p'$ as follows

$$\begin{aligned} \frac{1}{Q_p'} &= \frac{\sigma}{\omega_p \epsilon} + \frac{1}{\omega_p \mu} \frac{1}{1+j} \int Z_w \mathbf{H}_p^2 dS \\ &= \frac{\int \sigma [\mathbf{E}]_p \cdot [\mathbf{E}]_p^* dv}{\omega_p \int \epsilon [\mathbf{E}]_p \cdot [\mathbf{E}]_p^* dv} + \frac{\text{Re} \int Z_w [\mathbf{H}]_p \cdot [\mathbf{H}]_p^* dS}{\omega_p \int \mu [\mathbf{H}]_p \cdot [\mathbf{H}]_p^* dv} \\ &= \frac{\text{(power loss in medium)}}{2\omega_p \text{(average electric stored energy)}} \\ &\quad + \frac{\text{(power loss in wall)}}{2\omega_p \text{(average magnetic stored energy)}} \\ &= \frac{\text{(total power loss)}}{\omega_p \text{(total electromagnetic stored energy)}} \end{aligned}$$

In order to obtain (4.50), an assumption was made that only the p th mode was strongly excited in the cavity. However, if the mode is degenerate or if there are one or more ω_a 's close to ω_p , this assumption is no longer valid. In this case, several modes must be assumed, having the same order of magnitude, and whose resonant frequencies are close to the frequency under consideration. The right-hand side of (4.42) becomes a linear combination of the contributions from all of these modes, and corresponding to (4.45), we obtain simultaneous equations from which the expansion coefficients can be calculated. This situation is similar to the degenerate case in Section 3.8 in which the effect of waveguide wall losses is discussed.

4.3 Equivalent Circuits

In the previous section, the input admittance of a cavity with one opening was calculated and given by (4.54). If $\sigma = 0$, all the $(1/Q_a)$'s vanish and Y_p must become pure imaginary. In practice, however, a finite conductance may appear due to wall losses. This follows even though the conductance of each term in Y_p attributable to the wall losses is negligible compared to the susceptance of the same term. When the summation of an infinite number of terms is carried out, the conductance components add to a small, but finite, contribution while most of the susceptance components cancel each other. The effect of wall losses cannot be calculated by the simple approximate method we have used. However, the following discussion will not be affected since no assumption will be made as to the origin of the real part of Y_p .

For simplicity, let us assume that Y_p is a constant since it varies slowly with ω in the vicinity of ω_p , the frequency of interest. The inverse of the first term on the right-hand side of (4.54) expresses the impedance due to the resonant mode. Since the real part remains constant as ω varies, the locus of the impedance on the Smith chart must therefore be a constant resistance circle. The admittance locus, i.e., the inverse of the impedance, is a circle tangent to the periphery of the Smith chart where the reflection coefficient $r = -1$ as shown in Fig. 4.4(a). Adding Y_p to the circle, it follows from the property of bilinear transformations that the locus of Y becomes another circle as shown in Fig. 4.4(b). Let θ_p be the angle between the zero-susceptance line and the straight line passing through Y_p and the center of the Smith chart. Furthermore, let ω_0 be the angular frequency corresponding to the intersection of the admittance locus with this straight line as shown in Fig. 4.4(b). If the reference plane S_0 , at which the input admittance is

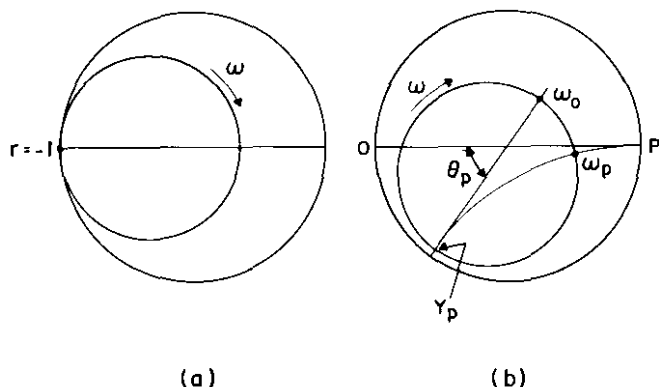


Fig. 4.4. Construction of the locus Y on a Smith chart.

calculated, is shifted toward the generator by

$$L_p = (\theta_p/4\pi) \lambda_g$$

the locus as a whole rotates clockwise by θ_p around the center of the Smith chart. As a result, the input admittance locus from the new reference plane may look like the circle shown in Fig. 4.5. Strictly speaking, the wavelength λ_g in the guide varies with ω . However, λ_g is assumed to be constant since we are considering a narrow range of ω in the vicinity of ω_p .

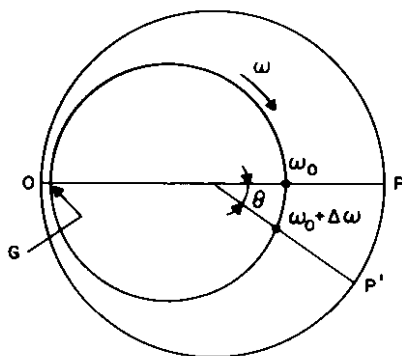


Fig. 4.5. Input admittance at the new reference plane S_0 .

With the new S_0 plane, the region V of the cavity is no longer the same as before, and the resonant frequencies together with the eigenfunctions are different from the original ones. However, the general expression for the input admittance should remain unchanged and is given by (4.54). As we can easily see from Fig. 4.5, ω_p' and Y_p for the new cavity must be given by ω_0 and a small conductance G , respectively. Thus, the normalized input admittance becomes

$$\frac{Y}{Y_0} = \frac{1/Q_{\text{ext}}}{j\{(\omega/\omega_0) - (\omega_0/\omega)\} + (1/Q_0)} + \frac{G}{Y_0} \quad (4.59)$$

where the external Q and the unloaded Q are defined through

$$1/Q_{\text{ext}} = I_p^2/\omega_p \mu Y_0, \quad 1/Q_0 = 1/Q_p'$$

for the new cavity, and Y_0 is the characteristic admittance of the propagating mode in the waveguide. The admittance Y can be considered as a parallel connection of a series resonant circuit and the conductance G . Comparing

the inverse of the resonant term

$$j \left\{ \left(\omega / \omega_0 \right) - \left(\omega_0 / \omega \right) \right\} \frac{Q_{\text{ext}}}{Y_0} + \frac{Q_{\text{ext}}}{Q_0 Y_0}$$

and an ordinary impedance expression for a series resonant circuit

$$j \{ \omega L - (1/\omega C) \} + R$$

we obtain an equivalent circuit of Y as shown in Fig. 4.6 where

$$L = Q_{\text{ext}}/\omega_0 Y_0, \quad C = Y_0/\omega_0 Q_{\text{ext}}, \quad R = Q_{\text{ext}}/Q_0 Y_0 \quad (4.60)$$

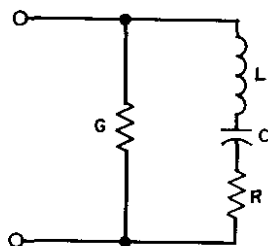


Fig. 4.6. Equivalent circuit of Y .

Let us consider the voltage at the original plane S_0 or at a reference plane distant $n\lambda_g/2$ from it, where n is a small integer. As a function of ω , this voltage changes proportionally to the straight-line distance from P to the point corresponding to ω on the locus in Fig. 4.4(b). Therefore, if ω is varied, while the incident voltage remains constant, the voltage at this plane should vary as illustrated by the curve in Fig. 4.7(a). On the other hand, the voltage as a function of ω at the new S_0 plane, or $n\lambda_g/2$ away from it, should look like the solid curve in Fig. 4.7(b). Consequently, the new S_0 plane can easily be determined with a standing wave detector. Note that the new S_0 plane is a voltage maximum point at frequencies far from resonance. The dotted line in Fig. 4.7(b) shows the voltage at a reference plane $\lambda_g/4$ away from the new S_0 plane. This voltage is proportional to the straight-line distance from the point 0 to the point corresponding to ω on the locus in Fig. 4.5. When the center of the Smith chart is inside the locus, the solid and dotted curves overlap each other as shown in Fig. 4.7(b); however, if the center is outside, the peak of the dotted curve becomes lower than the valley of the solid curve and no overlapping takes place. For a reason which will become clear after the discussion of power, the cavity is said to be overcoupled when

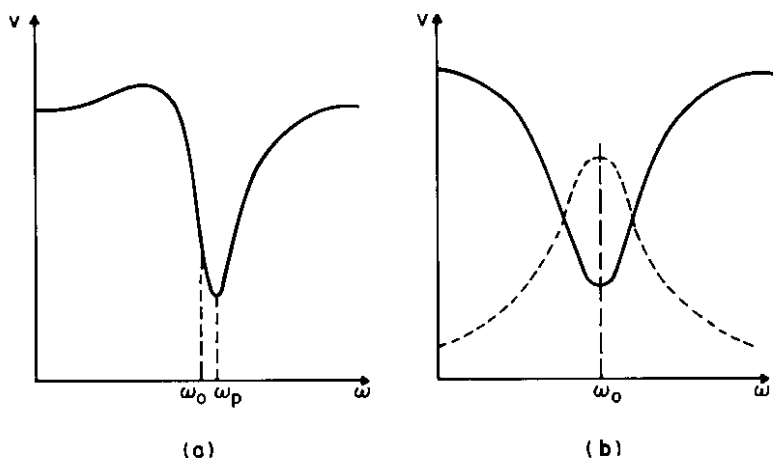


Fig. 4.7. Voltage versus ω at the reference planes.

the locus includes the center of the Smith chart, and undercoupled if it does not.

We next consider the position of the voltage minimum as a function of ω . In the case of overcoupling, referring to Fig. 4.5, the straight-line distance from P to the point corresponding to $\omega_0 + \Delta\omega$ on the locus is minimized when the locus as a whole is rotated counterclockwise around the center of the Smith chart through the angle θ . The counterclockwise rotation of the locus through θ means that the reference plane is shifted toward the load, a distance $L = (\theta/4\pi) \lambda_g$, and the length being minimum indicates that the voltage at this new reference plane is minimum. Thus, the position of voltage minimum at $\omega = \omega_0 + \Delta\omega$ is located $L = (\theta/4\pi) \lambda_g$ toward the load from the voltage minimum point for ω_0 . In the above procedure, instead of rotating the locus, P can be rotated in the opposite direction through θ to P' . The straight-line distance from P' to the point corresponding to $\omega_0 + \Delta\omega$ on the locus is then minimized indicating that the voltage minimum point for $\omega_0 + \Delta\omega$ shifts $L = (\theta/4\pi) \lambda_g$ toward the load from the minimum point for ω_0 . Repeating a similar procedure for each ω in the vicinity of ω_0 , the position of voltage minimum versus ω can be plotted as illustrated in Fig. 4.8(a). In the case of undercoupling, the locus does not enclose the center of the Smith chart, and the position of voltage minimum remains in the vicinity of the point O , shown in Fig. 4.5. The position of voltage minimum then changes with ω , as shown in Fig. 4.8(b), which is quite different from Fig. 4.8(a).

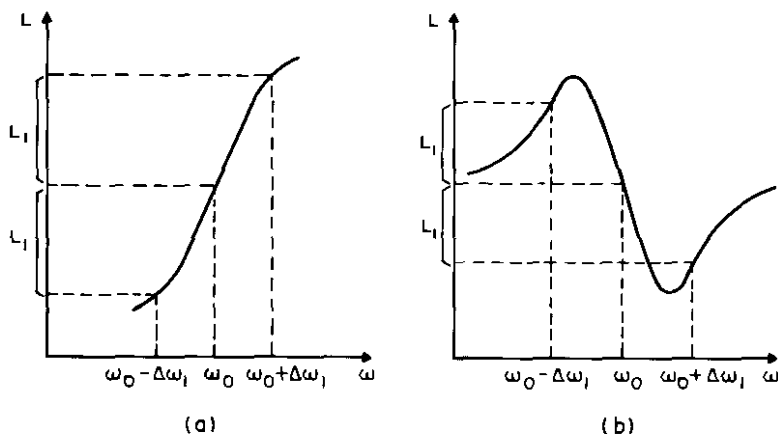


Fig. 4.8. Voltage minimum point versus ω (a) overcoupling; (b) undercoupling.

The value of G/Y_0 can be obtained from the standing wave ratio at frequencies sufficiently away from resonance (ω_0) that the first term on the right-hand side of (4.59) can be neglected compared to the second term. The frequencies should not be so far away that the value of Y_p is affected by the resonances of other modes. On the other hand, the value of $Q_0/Q_{ext} + G/Y_0$ can be obtained from the standing waves ratio at ω_0 . From these two values, we calculate

$$\frac{Y_1}{Y_0} = \frac{Q_0}{Q_{ext}(1 \pm j)} + \frac{G}{Y_0} \quad (4.61)$$

and obtain the straight-line distances from P or O to Y_1/Y_0 on the Smith chart, and hence the relative voltages at frequencies $\omega_1 = \omega_0 \pm \Delta\omega_1$ corresponding to Y_1/Y_0 . The frequency $\Delta\omega_1$ can be determined by measuring the voltage versus ω at an appropriate reference plane as shown in Fig. 4.7(b) and finding the frequencies at which the relative voltage becomes the value obtained above. Similarly, if θ , corresponding to Y_1/Y_0 is obtained on the Smith chart, the same frequencies $\omega_0 \pm \Delta\omega_1$ can be determined by taking two points distant $L_1 = (\theta/4\pi) \lambda_g$ from the point of symmetry on the measured curve of the voltage minimum point versus ω as shown in Fig. 4.8(a) or (b). Once $\Delta\omega_1$ is obtained, Q_0 can be calculated from

$$Q_0 = (\omega_0/2 \Delta\omega_1) \quad (4.62)$$

since a comparison of (4.59) and (4.61) shows that

$$1/Q_0 = |(\omega_1/\omega_0) - (\omega_0/\omega_1)| \simeq (2 \Delta\omega_1/\omega_0)$$

Furthermore, Q_{ext} can be determined from the measured values of G/Y_0 and $Q_0/Q_{\text{ext}} + G/Y_0$. Thus, all the parameters necessary to determine the equivalent circuit shown in Fig. 4.6 are experimentally obtainable.

Neglecting G , which is usually small, and assuming the generator admittance is equal to the characteristic admittance of the waveguide propagating mode, the equivalent circuit including the generator admittance is represented by Fig. 4.9. The Q of the circuit on the right-hand side of S_0 is given by $Q_0 = \omega_0 L/R$. The Q of the whole circuit including Y_0 is called the loaded Q and is indicated by Q_L ; and $1/Q_L$ is calculated as

$$\frac{1}{Q_L} = \frac{R + (1/Y_0)}{\omega_0 L} = \frac{1}{Q_0} + \frac{1}{Q_{\text{ext}}} \quad (4.63)$$

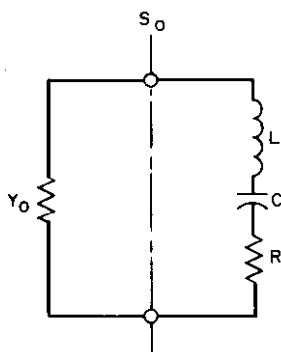


Fig. 4.9. Equivalent circuit of a cavity including the generator admittance.

Consequently, Q_{ext} can be interpreted as the Q corresponding to the power loss into Y_0 through the S_0 plane. If Q_{ext} is smaller than Q_0 , the energy escaping to the outside through the S_0 plane is larger than the energy consumed inside the cavity, and hence the cavity is overcoupled. On the other hand, if Q_{ext} is larger than Q_0 , the cavity is undercoupled to the outside. The center of the Smith chart is inside or outside the locus depending on whether the cavity is over- or undercoupled, respectively. If the effect of G is taken into account, the situation is different. Generally speaking, however, the cavity is said to be overcoupled whenever the locus includes the Smith chart center and is otherwise undercoupled. These two different states can easily be distinguished by performing an experiment similar to the one leading to Fig. 4.7(b) or Fig. 4.8.

Let us next discuss the equivalent circuit of a two-port resonant cavity.

To simplify the discussion, we restrict ourselves to the case in which the coupling between the input and output is due to one single resonant mode. In this case, the contribution to Y_{12} comes from the first term on the right-hand side of (4.57), and other terms become zero. Furthermore, if the S_0 plane in each waveguide is shifted to a voltage maximum point for frequencies far from resonance, as we did with the one-port cavity, Y_{11} as well as Y_{22} is represented by a resonant term plus a small conductance. We neglect the effect of the small conductances in Y_{11} and Y_{22} . Equation (4.56) then becomes

$$\begin{aligned} I_1 &= \frac{1}{\omega_0 \mu j \{(\omega/\omega_0) - (\omega_0/\omega)\} + (1/Q_0)} \left(\frac{I_{01}}{I_{02}} V_1 + V_2 \right) \\ I_2 &= \frac{1}{\omega_0 \mu j \{(\omega/\omega_0) - (\omega_0/\omega)\} + (1/Q_0)} \left(V_1 + \frac{I_{02}}{I_{01}} V_2 \right) \end{aligned} \quad (4.64)$$

If port II is short circuited, for example, by setting $V_2 = 0$, the input impedance of the resonant cavity from port I becomes a series resonant circuit. Similarly, if port I is short circuited, the cavity appears to port II as a series resonant circuit with the same resonant frequency as before, but the magnitude of the impedance is different. Therefore, an equivalent circuit of the cavity may look like Fig. 4.10, where a series resonant circuit Z is connected to the input and output circuits through transformers.

Let us study the relation between the currents and voltages in Fig. 4.10. The current I_0 flowing through Z is driven by the difference of electromotive forces $n_1 V_1$ and $n_2 V_2$, and we have

$$I_0 = Z^{-1} (n_1 V_1 - n_2 V_2)$$

From this, I_1 is given by

$$I_1 = n_1 I_0 = Z^{-1} \{n_1^2 V_1 - n_1 n_2 V_2\}$$

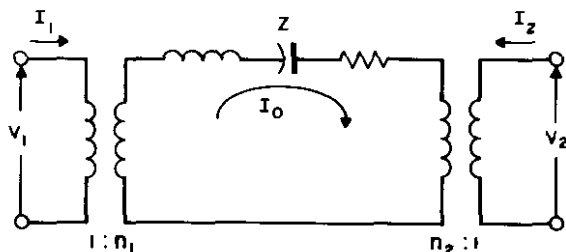


Fig. 4.10. Equivalent circuit of a two-port resonant cavity.

Similarly, I_2 becomes

$$I_2 = -n_2 I_0 = Z^{-1} \{-n_1 n_2 V_1 + n_2^2 V_2\}$$

Comparing these equations with (4.64), if the equivalent circuit is to represent the two-port cavity under consideration, we obtain

$$Z = \left\{ j \left(\frac{\omega}{\omega_0} - \frac{\omega_0}{\omega} \right) + \frac{1}{Q_0} \right\} \frac{\omega_0 \mu}{I_{01} I_{02}}$$

$$n_1 = (I_{01}/I_{02})^{1/2}, \quad n_2 = -(I_{02}/I_{01})^{1/2}$$

and Z is a series resonant circuit of L , C , and R given by

$$L = \mu/I_{01} I_{02}, \quad C = 1/\omega_0^2 L, \quad R = \omega_0 L/Q_0 \quad (4.65)$$

respectively.

Now, we are in a position to calculate the output power as a function of ω using this equivalent circuit. For simplicity, both generator and load are assumed to be matched to the waveguide characteristic impedances. The equivalent circuit, including the generator and load, is therefore given by Fig. 4.11(a) which becomes Fig. 4.11(b) when the transformers are eliminated. From Fig. 4.11(b), the output power is calculated to be

$$P_L(\omega) = n_2^2 R_L |I_0|^2$$

$$= \frac{n_2^2 R_L |n_1 E_g|^2}{|n_1^2 R_g + n_2^2 R_L + R + j\omega_0 L \{(\omega/\omega_0) - (\omega_0/\omega)\}|^2} \quad (4.66)$$

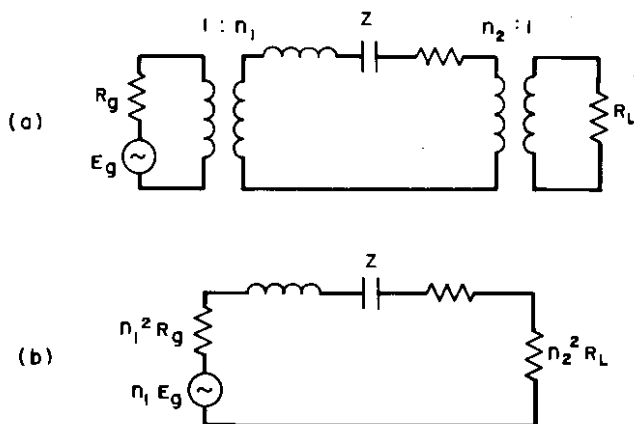


Fig. 4.11. Equivalent circuit of the two-port cavity including the generator and load.

If $Q_{\text{ext},1}$ and $Q_{\text{ext},2}$ are defined by

$$1/Q_{\text{ext},1} = n_1^2 R_g / \omega_0 L, \quad 1/Q_{\text{ext},2} = n_2^2 R_L / \omega_0 L \quad (4.67)$$

then, these external Q 's have a physical meaning similar to the external Q of one-port resonant cavities; i.e., $Q_{\text{ext},1}$ and $Q_{\text{ext},2}$ are the Q 's due to the power losses through ports 1 and 2, respectively. The loaded Q is then calculated from

$$(1/Q_L) = (1/Q_{\text{ext},1}) + (1/Q_{\text{ext},2}) + (1/Q_0) \quad (4.68)$$

In terms of Q_L , (4.66) can be rewritten in the form

$$P_L(\omega) = \frac{n_1^2 n_2^2 R_L E_g^2 Q_L^2}{(\omega_0 L)^2 |1 + jQ_L \{(\omega/\omega_0) - (\omega_0/\omega)\}|^2} \simeq \frac{P_L(\omega_0)}{1 + Q_L^2 (2\Delta\omega/\omega_0)^2} \quad (4.69)$$

where $\Delta\omega = \omega - \omega_0$. The output power as a function of ω becomes a typical bell shape as shown in Fig. 4.12. The difference of the frequencies at which the output power decreases to the one-half of the maximum is given by $2\Delta\omega = \omega_0/Q_L$. This relation is often used for the measurement of Q_L .

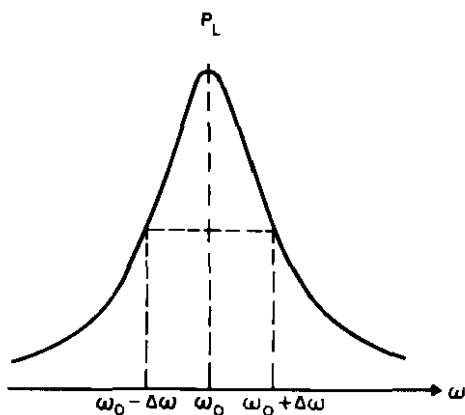


Fig. 4.12. Output power versus ω of the two-port cavity.

1.4 Perturbation of Boundaries

Suppose the cavity wall is slightly deformed from S' to S'' as shown in Fig. 4.13, where the deformed surface is indicated by ΔS . The volume of the cavity is now $V - \Delta V$, where ΔV is the volume of the deformed part. The eigenfunctions as well as the eigenvalues will now all be different from

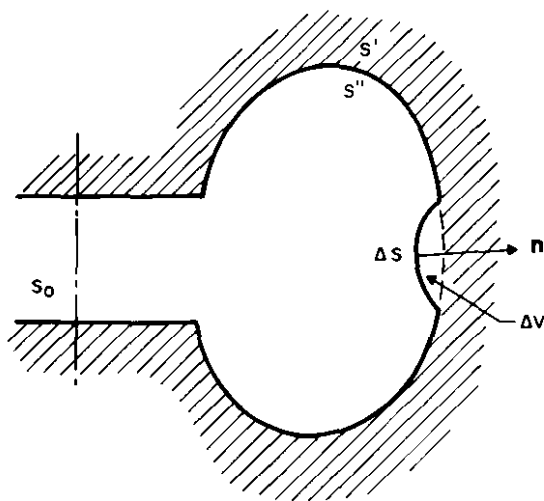


Fig. 4.13. A small perturbation of the wall of a cavity.

the original ones. As a result, I_{p1}^2 , ω_p' , and Y_p in (4.54) are expected to change. However, since the denominator of the first term on the right-hand side almost cancels in the vicinity of ω_p' , the largest effect on the input admittance comes from the variation in ω_p' . For this reason, the variation of ω_p' due to the deformation is particularly worth studying. However, since ω_p' is directly proportional to ω_p and Q_{Sp} is expected to remain almost the same, we have only to investigate the variation of ω_p , or k_p^2 which is equal to $\omega_p^2 \epsilon \mu$, in order to study the ω_p' variation.

Let us use the variational expression for k^2 to investigate the variation of k_p^2 . Noting that the expression gives an approximation for the eigenvalue one order higher than that of the trial function used, we use the original eigenfunction $\mathbf{E}_p$ as a trial function. Then, $k^2(\mathbf{E}_p)$ is given by

$$\begin{aligned}
 k^2(\mathbf{E}_p) &= \frac{\int_{V-\Delta V} (\nabla \times \mathbf{E}_p)^2 dv - 2 \int_{S''+S_0} \mathbf{n} \times \mathbf{E}_p \cdot \nabla \times \mathbf{E}_p dS}{\int_{V-\Delta V} \mathbf{E}_p^2 dv} \\
 &= \frac{\int_V (\nabla \times \mathbf{E}_p)^2 dv - \int_{\Delta V} (\nabla \times \mathbf{E}_p)^2 dv - 2 \int_{S''+S_0} \mathbf{n} \times \mathbf{E}_p \cdot \nabla \times \mathbf{E}_p dS}{\int_V \mathbf{E}_p^2 dv - \int_{\Delta V} \mathbf{E}_p^2 dv}
 \end{aligned}$$

$$\begin{aligned} &\simeq k_p^2 - \int_{\Delta V} (\nabla \times \mathbf{E}_p)^2 dv + k_p^2 \int_{\Delta V} \mathbf{E}_p^2 dv \\ &\quad - 2 \int_{S''+S_0} \mathbf{n} \times \mathbf{E}_p \cdot \nabla \times \mathbf{E}_p dS \end{aligned}$$

where $(\nabla \cdot \mathbf{E}_p = 0)$ and the normalization condition of $\mathbf{E}_p$ are used. The last surface integral is equal to the integral over ΔS since $\mathbf{n} \times \mathbf{E}_p = 0$ elsewhere. However, using Gauss's theorem and $\mathbf{n} \times \mathbf{E}_p = 0$ on S' , we have

$$\begin{aligned} \int_{\Delta S} \mathbf{n} \times \mathbf{E}_p \cdot \nabla \times \mathbf{E}_p dS &= - \int_{\Delta V} \nabla \cdot \mathbf{E}_p \times \nabla \times \mathbf{E}_p dv \\ &= - \int_{\Delta V} \{(\nabla \times \mathbf{E}_p)^2 - \mathbf{E}_p \cdot \nabla \times \nabla \times \mathbf{E}_p\} dv \\ &= - \int_{\Delta V} (\nabla \times \mathbf{E}_p)^2 dv + \int_{\Delta V} k_p^2 \mathbf{E}_p^2 dv \end{aligned}$$

Combining the above two equations, we have

$$k^2(\mathbf{E}_p) \simeq k_p^2 + \int_{\Delta V} (\nabla \times \mathbf{E}_p)^2 dv - k_p^2 \int_{\Delta V} \mathbf{E}_p^2 dv$$

It follows from (4.19) that this is equivalent to

$$k^2(\mathbf{E}_p) = k_p^2 \left\{ 1 + \int_{\Delta V} (\mathbf{H}_p^2 - \mathbf{E}_p^2) dv \right\} \quad (4.70)$$

Equation (4.70) gives the approximate eigenvalue for the deformed cavity. Writing the variation in the resonant frequency by $\Delta\omega_p$, we obtain

$$\frac{\Delta\omega_p}{\omega_p} \simeq \frac{k^2(\mathbf{E}_p) - k_p^2}{2k_p^2} = \frac{1}{2} \int_{\Delta V} (\mathbf{H}_p^2 - \mathbf{E}_p^2) dv \quad (4.71)$$

This can also be considered to give $\Delta\omega_p'/\omega_p'$. From (4.71) we see that the resonant frequency increases or decreases depending on whether the magnetic-stored energy is larger or smaller than the electric-stored energy in the region removed from the cavity by the deformation. In other words, if the wall is pushed in where the electric or magnetic field is concentrated, the resonant frequency will be lowered or raised, respectively. This corresponds to the behavior of an *LC* resonant circuit where the resonant frequency goes down if we bring the electrodes of the capacitor closer together; whereas, if we insert a piece of metal inside the coil, the inductance decreases and the resonant frequency goes up.

In the above discussion, when $\mathbf{E}_p$ was used as a trial function, it was tacitly assumed that $\mathbf{E}_p$ was not a degenerate eigenfunction. If it is, a linear combination of all the eigenfunctions which are degenerate to $\mathbf{E}_p$ must be used in place of $\mathbf{E}$ in the variational expression, and the coefficients have to be determined so as to make $k^2(\mathbf{E})$ stationary. Each stationary value of $k^2(\mathbf{E})$ gives an eigenvalue from which the resonant frequency can be calculated. The situation is somewhat similar to the degenerate case discussed in Section 3.8.

4.5 Cavities with Inhomogeneous Media

In this section, let us discuss briefly how to obtain the equivalent circuit of resonant cavities with inhomogeneous media. Since the two sets of eigenfunctions obtained in Section 4.2 are complete, it is possible to expand the expressions for electric and magnetic fields in a cavity with an inhomogeneous medium in terms of these sets of eigenfunctions. If this is done, however, it becomes impossible to assume that only one or a small number of modes are excited strongly, and that others are negligible. In other words, the convergence of the series becomes so poor that no useful information can be obtained. Consequently, some other sets of functions have to be found which are suitable for the expansion of electromagnetic fields in such cavities. Fortunately, the following eigenvalue problems give suitable sets of eigenfunctions provided that $\epsilon_r = \epsilon/\epsilon_0$ and $\mu_r = \mu/\mu_0$ remain positive and finite everywhere in the cavity.

$$\begin{aligned} \nabla \times \mu_r^{-1} \nabla \times \mathbf{E} - \epsilon_r \nabla \nabla \cdot \epsilon_r \mathbf{E} - k^2 \epsilon_r \mathbf{E} &= 0 & (\text{in } V) \\ \mathbf{n} \times \mathbf{E} &= 0, \quad \nabla \cdot \epsilon_r \mathbf{E} = 0 & (\text{on } S) \end{aligned} \quad (4.72)$$

$$\begin{aligned} \nabla \times \epsilon_r^{-1} \nabla \times \mathbf{H} - \mu_r \nabla \nabla \cdot \mu_r \mathbf{H} - k^2 \mu_r \mathbf{H} &= 0 & (\text{in } V) \\ \mathbf{n} \cdot \mu_r \mathbf{H} &= 0, \quad \mathbf{n} \times \epsilon_r^{-1} \nabla \times \mathbf{H} = 0 & (\text{on } S) \end{aligned} \quad (4.73)$$

Under the assumption that ϵ_r and μ_r are positive and finite, k^2 becomes real and nonnegative, and the eigenfunctions can be chosen to be real, without loss of generality. The variational expressions for the k^2 's are given by

$$k^2(\mathbf{E}) = \frac{\int \mu_r^{-1} (\nabla \times \mathbf{E})^2 dv + \int (\nabla \cdot \epsilon_r \mathbf{E})^2 dv - 2 \int \mathbf{n} \times \mathbf{E} \cdot \mu_r^{-1} \nabla \times \mathbf{E} dS}{\int \epsilon_r \mathbf{E}^2 dv} \quad (4.74)$$

$$k^2(\mathbf{H}) = \frac{\int \varepsilon_r^{-1} (\mathbf{V} \times \mathbf{H})^2 dv + \int (\mathbf{V} \cdot \mu_r \mathbf{H})^2 dv - 2 \int \mathbf{n} \cdot \mu_r \mathbf{H} \mathbf{V} \cdot \mu_r \mathbf{H} dS}{\int \mu_r \mathbf{H}^2 dv} \quad (4.75)$$

Using these expressions, the eigenfunctions can be obtained, at least conceptionally, one by one to satisfy the orthonormal conditions

$$\int \varepsilon_r \mathbf{E}_n \cdot \mathbf{E}_m dv = \begin{cases} 0 & (n \neq m) \\ 1 & (n = m) \end{cases} \quad (4.76)$$

$$\int \mu_r \mathbf{H}_n \cdot \mathbf{H}_m dv = \begin{cases} 0 & (n \neq m) \\ 1 & (n = m) \end{cases} \quad (4.77)$$

In this process, every $\mathbf{E}_n$ can be made to belong to one of three groups

$$\begin{aligned} \text{I. } & \mathbf{V} \times \mathbf{E}_n = 0, & \mathbf{V} \cdot \varepsilon_r \mathbf{E}_n = 0 \\ \text{II. } & \mathbf{V} \times \mathbf{E}_n \neq 0, & \mathbf{V} \cdot \varepsilon_r \mathbf{E}_n = 0 \\ \text{III. } & \mathbf{V} \times \mathbf{E}_n = 0, & \mathbf{V} \cdot \varepsilon_r \mathbf{E}_n \neq 0 \end{aligned}$$

Similarly, we have for $\mathbf{H}_m$

$$\begin{aligned} \text{I. } & \mathbf{V} \times \mathbf{H}_m = 0, & \mathbf{V} \cdot \mu_r \mathbf{H}_m = 0 \\ \text{II. } & \mathbf{V} \times \mathbf{H}_m \neq 0, & \mathbf{V} \cdot \mu_r \mathbf{H}_m = 0 \\ \text{III. } & \mathbf{V} \times \mathbf{H}_m = 0, & \mathbf{V} \cdot \mu_r \mathbf{H}_m \neq 0 \end{aligned}$$

Furthermore, each function belonging to group II can be selected so as to make a pair, $\mathbf{E}_a$ and $\mathbf{H}_a$, satisfying

$$\mathbf{V} \times \mathbf{E}_a = k_a \mu_r \mathbf{H}_a, \quad \mathbf{V} \times \mathbf{H}_a = k_a \varepsilon_r \mathbf{E}_a \quad (4.78)$$

If Greek subscripts are used for the functions belonging to groups I and III, we have

$$\mathbf{V} \times \mathbf{E}_\nu = 0, \quad \mathbf{V} \times \mathbf{H}_\lambda = 0 \quad (4.79)$$

Comparing the magnitudes of the eigenvalues for cavities with homogeneous and inhomogeneous media, the infinite growth of the eigenvalues of the present problems can be proved as shown in Appendix I. Therefore, each set of the orthonormal functions obtained above is complete. An arbitrary vector function $\mathbf{F}$ which is piecewise-continuous and square-integrable over V can be expanded in terms of either set of these functions; i.e.,

$$\mathbf{F} = \sum_{n=1}^{\infty} A_n \mathbf{E}_n, \quad \mathbf{F} = \sum_{m=1}^{\infty} B_m \mathbf{H}_m \quad (4.80)$$

where

$$A_n = \int \epsilon_r \mathbf{F} \cdot \mathbf{E}_n dv, \quad B_m = \int \mu_r \mathbf{F} \cdot \mathbf{H}_m dv \quad (4.81)$$

For the analysis of the cavity, the quantities appearing in Maxwell's equations are to be expanded in terms of these sets of functions. We use the $\mathbf{E}_n$'s and the $\mathbf{H}_m$'s to expand $\mathbf{E}$ and $\mathbf{H}$, respectively. However, if we try to expand $\nabla \times \mathbf{E}$ and $\nabla \times \mathbf{H}$, we shall find that the convergences are poor since they do not resemble $\mathbf{H}_n$ and $\mathbf{E}_n$, respectively. Therefore, we expand $\mu_r^{-1} \nabla \times \mathbf{E}$ and $\epsilon_r^{-1} \nabla \times \mathbf{H}$ in terms of the $\mathbf{H}_m$'s and the $\mathbf{E}_n$'s, instead of $\nabla \times \mathbf{E}$ and $\nabla \times \mathbf{H}$. The remainder of the discussion then becomes almost identical to that for cavities with homogeneous media. The expression for the input admittance remains the same as (4.54), and the discussions in Section 4.3 apply equally well, without modification, to the present inhomogeneous case.

PROBLEMS

- 4.1 Try to prove the completeness of the eigenfunctions defined by the eigenvalue problem: $\nabla \times \nabla \times \mathbf{E} - k^2 \mathbf{E} = 0$ (in V) and $\mathbf{n} \times \mathbf{E} = 0$ (on S). Point out where the discussion fails.
- 4.2 Calculate the change in the resonant frequency of an LC resonant circuit due to the insertion of a dielectric slab between the electrodes of the capacitor as shown in Fig. 4.14. Use two different methods: (1) Calculating the change of the capacitance; (2) Using the variational expression (4.74).

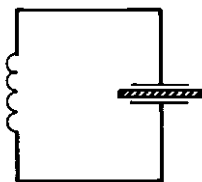


Fig. 4.14. LC resonator with a dielectric slab between the capacitance electrodes.

- 4.3 Suppose that a piece of waveguide an integer multiple of a half-wavelength long is short-circuited at both ends to form a cavity. Show that the attenuation constant of the waveguide can be calculated from the Q_{sp} of the cavity through $\alpha = (\pi \lambda_g / Q_{sp} \lambda^2)$ where the wall losses in the end plates are neglected.
- 4.4 Considering a length L of a rectangular waveguide short circuited at one end as a cavity, obtain the $\mathbf{H}_a$'s and $\mathbf{H}_\lambda$'s, explicitly.

- 4.5 Using the above expressions and neglecting the wall losses, calculate the input admittance of the rectangular waveguide (TE_{10} mode) at the reference plane distant L from the short-circuited end.
- 4.6 Prove that the input admittance in Problem 4.5 is equal to that obtainable by treating the waveguide as a transmission line. (*Hint*: A well-known formula

$$\cot \theta = \frac{1}{\theta} \left\{ 1 + \sum_{n=1}^{\infty} \frac{2}{1 - (n\pi/\theta)^2} \right\}$$

may be helpful to equate the two expressions.)

- 4.7 Show that the solutions of (4.8) under the boundary conditions

$$\begin{aligned} \mathbf{n} \times \nabla \times \mathbf{E} &= 0, & \mathbf{n} \cdot \mathbf{E} &= 0 & (\text{on } S_0) \\ \mathbf{n} \times \mathbf{E} &= 0, & \nabla \cdot \mathbf{E} &= 0 & (\text{on } S') \end{aligned}$$

form a complete set of orthonormal functions. Also show that the solutions of (4.14) under the boundary conditions

$$\begin{aligned} \mathbf{n} \times \mathbf{H} &= 0, & \nabla \cdot \mathbf{H} &= 0 & (\text{on } S_0) \\ \mathbf{n} \times \nabla \times \mathbf{H} &= 0, & \mathbf{n} \cdot \mathbf{H} &= 0 & (\text{on } S') \end{aligned}$$

form another complete set of orthonormal functions. The boundary conditions on S_0 correspond to the open-circuited condition.

- 4.8 Using the above sets of eigenfunctions, obtain the input impedance of the cavity at the reference plane S_0 .
- 4.9 Calculate the input admittance of a cavity with twofold degenerate resonant modes.
- 4.10 Use a perturbation method, similar to the one used in Section 3.8, to obtain (4.71).
- 4.11 Carry out the discussion of boundary perturbation assuming the twofold degeneracy of the resonant modes.
- 4.12 Complete the discussion for obtaining an equivalent circuit representing cavities with inhomogeneous media.

CHAPTER 5

MATRICES AND WAVEGUIDE JUNCTIONS

In Chapter 2, some fundamental properties of vectors were reviewed. It was shown that a single letter could represent three components of a vector, and relations between the components of various vectors were conveniently studied using certain rules of mathematical vector operations. The introduction of vectors provided a space saving technique for describing these relations, and also reduced the tremendous mental effort otherwise necessary to handle three times as many variables representing vector components.

In much the same way, a matrix represents several quantities in a prescribed manner. Using certain operating rules between matrices, it becomes possible to describe complex relations between many quantities in an orderly fashion. Thus, by the introduction of matrices, the description of the input and output relations of complex waveguide junctions becomes simpler, and the understanding of their behavior becomes easier. As a result, several useful theorems for treating microwave circuits, which otherwise would be difficult to find, now become readily obtainable.

In this chapter, we shall first present matrix analysis which has proved useful in many branches of applied physics, including the theory of microwave circuits. Then, using matrix notations, the reciprocal theorem, lossless conditions, and frequency characteristics of lossless reciprocal circuits will be described. Also, the properties of symmetrical Y junctions and circulators will be discussed in detail.

5.1 Matrices

A matrix $\mathbf{A}$ of order $m \times n$ is a particular collection of $m \times n$ quantities.

They are generally arranged in m rows and n columns as follows:

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1n} \\ a_{21} & a_{22} & \cdots & a_{2n} \\ \vdots & \vdots & \vdots & \vdots \\ a_{m1} & a_{m2} & \cdots & a_{mn} \end{bmatrix} \quad (5.1)$$

and a_{ij} is called the ij component of $\mathbf{A}$. Two matrices $\mathbf{A}$ and $\mathbf{B}$ are said to be equal when, and only when, all the corresponding components exist and are equal to each other; i.e., $a_{ij} = b_{ij}$. Therefore, when they are equal, the matrices are of the same order. The addition of two matrices $\mathbf{A}$ and $\mathbf{B}$ is defined as a matrix $\mathbf{C}$ with its ij component c_{ij} being $a_{ij} + b_{ij}$. Addition is only defined between matrices of the same order. Since $a_{ij} + b_{ij} = b_{ij} + a_{ij}$, we have

$$\mathbf{C} = \mathbf{A} + \mathbf{B} = \mathbf{B} + \mathbf{A} \quad (5.2)$$

That is, addition is commutative; the order of addition is interchangeable without changing the result.

The product of matrices $\mathbf{A}$ of order $m \times l$ and $\mathbf{B}$ of order $l \times n$ is defined as a matrix $\mathbf{C}$ of order $m \times n$ with its ij component being

$$c_{ij} = \sum_k a_{ik} b_{kj}$$

Multiplication between matrices is defined only when the number of columns in the first matrix is equal to the number of rows in the second matrix. When matrices of order 2×3 and 3×2 are multiplied, we obtain a matrix of order 2×2 ,

$$\begin{bmatrix} a_{11} & a_{12} & a_{13} \\ -a_{21} & -a_{22} & -a_{23} \end{bmatrix} \begin{bmatrix} \downarrow & & \\ b_{11} & b_{12} & \\ \downarrow & & \\ b_{21} & b_{22} & \\ \downarrow & & \\ b_{31} & b_{32} & \\ \downarrow & & \end{bmatrix} = \begin{bmatrix} c_{11} & c_{12} \\ \textcircled{c_{21}} & c_{22} \end{bmatrix}$$

Here, for example, the elements of the second row in $\mathbf{A}$ are multiplied by corresponding elements of the first column in $\mathbf{B}$, and added together to yield the 21 component c_{21} of $\mathbf{C}$; $a_{21}b_{11} + a_{22}b_{21} + a_{23}b_{31} = c_{21}$. From the definitions of addition and multiplication, it follows that

$$\mathbf{A}(\mathbf{B} + \mathbf{C}) = \mathbf{AB} + \mathbf{AC} \quad (5.3)$$

This is because the ij component of the matrix on the left-hand side is given by

$$\sum_k a_{ik}(b_{kj} + c_{kj}) = \sum_k a_{ik}b_{kj} + \sum_k a_{ik}c_{kj}$$

which is exactly equal to the sum of the ij components of the matrices on the right-hand side. Similarly, when multiplying a matrix by the product of two matrices, we have

$$\mathbf{A}(\mathbf{BC}) = (\mathbf{AB})\mathbf{C} \quad (5.4)$$

where the parentheses indicate the first operation to be performed. It is worth noting that $\mathbf{AB}$ is not necessarily equal to $\mathbf{BA}$.

A zero matrix $\mathbf{0}$ is a matrix whose components are all equal to zero. It is different from zero which is usually indicated by 0. However, when 0 appears in matrix equations, it may mean a $\mathbf{0}$ matrix. The product of a $\mathbf{0}$ matrix and an ordinary matrix makes another $\mathbf{0}$ matrix:

$$\mathbf{A}\mathbf{0} = \mathbf{0}, \quad \mathbf{0}\mathbf{A} = \mathbf{0} \quad (5.5)$$

Although we use the same symbol $\mathbf{0}$ for all zero matrices, they may not be the same matrix; the numbers of rows and columns may be different for different $\mathbf{0}$'s in (5.5).

A matrix with the same number of rows and columns (m) is called a square matrix of order m ; it is really a matrix of order $m \times m$. A unit matrix $\mathbf{I}$ is a square matrix whose main diagonal components are each equal to unity while all the other components are zero. We have

$$\mathbf{IA} = \mathbf{A}, \quad \mathbf{AI} = \mathbf{A} \quad (5.6)$$

In the multiplication of a matrix with a constant, all the components of the matrix are multiplied by the constant. Thus, $c\mathbf{I}$ is a square matrix with each of its main diagonal components equal to c while the remainder are zero. It follows from this that

$$c\mathbf{A} = (c\mathbf{I})\mathbf{A} \quad (5.7)$$

The inverse matrix of a square matrix $\mathbf{A}$ is indicated by $\mathbf{A}^{-1}$, and it is defined as a matrix which satisfies the following relations:

$$\mathbf{AA}^{-1} = \mathbf{I}, \quad \mathbf{A}^{-1}\mathbf{A} = \mathbf{I} \quad (5.8)$$

From this definition, we see that if the inverse matrix of a matrix exists, it is unique.

Let $\det \mathbf{A}$ be the determinant with the same components as matrix $\mathbf{A}$,

and let A_{ij} be the cofactor of the ij component. Then, since

$$\begin{aligned}\det \mathbf{A} &= \sum_j a_{ij} A_{ij} = \sum_i a_{ij} A_{ij} \\ \sum_j a_{ij} A_{kj} &= 0 \quad (i \neq k) \\ \sum_i a_{ij} A_{ik} &= 0 \quad (j \neq k)\end{aligned}\tag{5.9}$$

the ji component of $\mathbf{A}^{-1}$ is given by A_{ij} divided by $\det \mathbf{A}$. This can be easily checked by substituting into (5.8). When $\det \mathbf{A} = 0$, $\mathbf{A}^{-1}$ does not exist. A square matrix $\mathbf{A}$ without $\mathbf{A}^{-1}$ is said to be singular; when $\mathbf{A}^{-1}$ exists, $\mathbf{A}$ is nonsingular. A nonsingular matrix is a square matrix by definition. The inverse matrix of the product of two nonsingular matrices $\mathbf{A}$ and $\mathbf{B}$ is given by

$$(\mathbf{AB})^{-1} = \mathbf{B}^{-1} \mathbf{A}^{-1}\tag{5.10}$$

This is because $\mathbf{B}^{-1} \mathbf{A}^{-1}$ satisfies the definition (5.8) of the inverse matrix as follows:

$$\mathbf{ABB}^{-1} \mathbf{A}^{-1} = \mathbf{AA}^{-1} = \mathbf{I}, \quad \mathbf{B}^{-1} \mathbf{A}^{-1} \mathbf{AB} = \mathbf{B}^{-1} \mathbf{B} = \mathbf{I}$$

Similarly, for the product of three matrices, we have

$$(\mathbf{ABC})^{-1} = \mathbf{C}^{-1} \mathbf{B}^{-1} \mathbf{A}^{-1}\tag{5.11}$$

Hence, when the inverse is taken, the order of product is reversed.

The transposed matrix $\mathbf{A}_t$ of a matrix $\mathbf{A}$ is the one with the rows and columns interchanged. For example,

$$\mathbf{A} = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \end{bmatrix}, \quad \mathbf{A}_t = \begin{bmatrix} a_{11} & a_{21} \\ a_{12} & a_{22} \\ a_{13} & a_{23} \end{bmatrix}$$

The ij component of the transposed matrix of a product $\mathbf{AB}$ is given by

$$\sum_k a_{jk} b_{ki} = \sum_k b_{ki} a_{jk}$$

However, since the right-hand side is equal to the ij component of $\mathbf{B}_t \mathbf{A}_t$, we have

$$(\mathbf{AB})_t = \mathbf{B}_t \mathbf{A}_t\tag{5.12}$$

The order of product is reversed for the transposed matrix. In the case of the product of three matrices, we have

$$(\mathbf{ABC})_t = \mathbf{C}_t \mathbf{B}_t \mathbf{A}_t\tag{5.13}$$

If the transposed matrix of A has all its components changed to their complex conjugates, the matrix thus obtained is called the adjoint matrix of A and is indicated by A^+ . Just as the transposed matrix reverses the order of a product, so does the adjoint matrix, i.e.,

$$(AB)^+ = B^+ A^+ \quad (5.14)$$

When A^+ is equal to A , A is called a self-adjoint matrix. If

$$A^+ = A^{-1} \quad (5.15)$$

then A is called a unitary matrix. The product of two unitary matrices is again a unitary matrix since

$$(AB)^+ = B^+ A^+ = B^{-1} A^{-1} = (AB)^{-1} \quad (5.16)$$

where (5.10) and (5.14) are used together with (5.15).

A constant times a vector in three dimensional space is a vector whose components are multiplied by the constant. The sum of vectors is defined to be a vector with each component representing the sum of the corresponding components in the vectors. Consequently, there is one-to-one correspondence between matrices of order 3×1 and vectors in the three-dimensional space. Extending this correspondence to multidimensional spaces, let us call a matrix of order $n \times 1$ a vector in n -dimensional space. The operation of multiplying a square matrix A of order n and a matrix x of order $n \times 1$ to obtain another matrix y of order $n \times 1$, can then be expressed as follows: Vector x is transformed into vector y by A . When A is a unitary matrix, we have

$$y^+ y = (Ax)^+ Ax = x^+ A^+ Ax = x^+ x \quad (5.17)$$

where use is made of (5.15). When the components of the vector x in three-dimensional space are all real, $x^+ x$ expresses the square of the magnitude of x . Extending this notion of magnitude to more general cases, a real quantity $x^+ x$ is considered as the square of the magnitude of vector x even when x is an n -dimensional vector with complex components. With this interpretation, (5.17) shows that the magnitude of a vector remains unchanged during the transformation by a unitary matrix. In other words, the magnitudes of vectors are invariant to unitary transformations.

Now suppose that

$$y = Ax \quad (5.18)$$

and let us study how A changes during a coordinate transformation. If T

is a nonsingular matrix which expresses the coordinate transformation, and if $\mathbf{x}'$ and $\mathbf{y}'$ are the results of the transformation, then

$$\mathbf{x}' = \mathbf{T}\mathbf{x}, \quad \mathbf{y}' = \mathbf{T}\mathbf{y} \quad (5.19)$$

Substituting into (5.18), we have

$$\mathbf{T}^{-1}\mathbf{y}' = \mathbf{A}\mathbf{T}^{-1}\mathbf{x}'$$

or equivalently

$$\mathbf{y}' = \mathbf{TAT}^{-1}\mathbf{x}'$$

Comparing this with

$$\mathbf{y}' = \mathbf{A}'\mathbf{x}'$$

we see that $\mathbf{A}$ is transformed into

$$\mathbf{A}' = \mathbf{TAT}^{-1} \quad (5.20)$$

The transformation expressed by (5.20) is called the similarity transformation by $\mathbf{T}$. This name comes from the fact that the form of a matrix equation is preserved during the transformation. For example, suppose

$$\mathbf{AB} + \mathbf{CDE} = \mathbf{F}$$

After the similarity transformation by $\mathbf{T}$, we have

$$\mathbf{TABT}^{-1} + \mathbf{TCDET}^{-1} = \mathbf{TFT}^{-1}$$

which can be rewritten in the form

$$\mathbf{TAT}^{-1}\mathbf{TB T}^{-1} + \mathbf{TCT}^{-1}\mathbf{TDT}^{-1}\mathbf{TET}^{-1} = \mathbf{TFT}^{-1}$$

or equivalently,

$$\mathbf{A}'\mathbf{B}' + \mathbf{C}'\mathbf{D}'\mathbf{E}' = \mathbf{F}'$$

This has the same form as the original equation. Thus, the forms of matrix equations are invariant to similarity transformations.

When a vector is transformed by $\mathbf{A}$, its direction generally changes, i.e., the ratios between components become different from the original values. However, there are vectors whose directions are invariant to a given transformation $\mathbf{A}$. Let us try to find such vectors. Since vector directions do not change by constant multiplications, we look for vectors which satisfy

$$\mathbf{Ax} = \lambda\mathbf{x} \quad (5.21)$$

or equivalently

$$(\mathbf{A} - \lambda\mathbf{I})\mathbf{x} = \mathbf{0}$$

where 0 on the right-hand side means a $\mathbf{0}$ matrix as we mentioned earlier. The above equation is equivalent to n simultaneous equations

$$\begin{aligned}(a_{11} - \lambda)x_1 + a_{12}x_2 + \cdots + a_{1n}x_n &= 0 \\ a_{21}x_1 + (a_{22} - \lambda)x_2 + \cdots + a_{2n}x_n &= 0 \\ &\vdots \\ a_{n1}x_1 + a_{n2}x_2 + \cdots + (a_{nn} - \lambda)x_n &= 0\end{aligned}$$

For nontrivial solutions to exist, the determinant of the coefficients must be equal to zero:

$$\det(\mathbf{A} - \lambda\mathbf{I}) = 0 \quad (5.22)$$

This is an n th degree algebraic equation for λ which will generally have n different roots. Corresponding to the n roots, n independent vectors $\mathbf{x}_i$ ($i = 1, 2, \dots, n$) satisfying (5.21) can be determined; the vector $\mathbf{x}_i$ is called an eigenvector, and the corresponding root λ_i is called the eigenvalue. We use this terminology because of the similarity between this and the eigenvalue problems discussed previously in connection with waveguides and cavities. In those problems, we essentially searched for vector functions which became k^2 times the original ones when certain differential operations were performed.

Let us define $f(\lambda)$ by

$$f(\lambda) = \det(\lambda\mathbf{I} - \mathbf{A}) \quad (5.23)$$

Then, (5.22) can be written as $f(\lambda) = 0$. The form of $f(\lambda)$ is invariant to the similarity transformation of $\mathbf{A}$ by $\mathbf{T}$; the proof is as follows. From (5.20), we have $\mathbf{T}^{-1}\mathbf{A}' = \mathbf{A}\mathbf{T}^{-1}$. Therefore, we obtain

$$\mathbf{T}^{-1}(\lambda\mathbf{I} - \mathbf{A}') = (\lambda\mathbf{I} - \mathbf{A})\mathbf{T}^{-1}$$

Taking the determinants of both sides, we have

$$\det \mathbf{T}^{-1} \det(\lambda\mathbf{I} - \mathbf{A}') = \det(\lambda\mathbf{I} - \mathbf{A}) \det \mathbf{T}^{-1}$$

Since $\mathbf{T}$ is nonsingular and $\det \mathbf{T}^{-1}$ is not equal to zero, the above expression is equivalent to

$$\det(\lambda\mathbf{I} - \mathbf{A}') = \det(\lambda\mathbf{I} - \mathbf{A}) \quad (5.24)$$

which is the result we wished to prove.

Let $\lambda_1, \lambda_2, \dots, \lambda_n$ be the eigenvalues. Then we have

$$f(\lambda) = (\lambda - \lambda_1)(\lambda - \lambda_2) \cdots (\lambda - \lambda_n) \quad (5.25)$$

Since $f(\lambda)$ is invariant, the eigenvalues are also invariant to similarity

transformations of $\mathbf{A}$. Furthermore, since $f(\lambda)$ is an n th-degree polynomial of λ , it can be rewritten in the form

$$f(\lambda) = \lambda^n - s_1 \lambda^{n-1} + \cdots + (-1)^n s_n \quad (5.26)$$

Because of the invariance of $f(\lambda)$, $s_1, s_2, \dots, s_n$ are all invariant to similarity transformations of $\mathbf{A}$. In particular, s_1 is given by

$$s_1 = a_{11} + a_{22} + \cdots + a_{nn} = \lambda_1 + \lambda_2 + \cdots + \lambda_n \quad (5.27)$$

It is called the trace of $\mathbf{A}$ and is indicated by $\text{tr } \mathbf{A}$. Setting λ is equal to zero, s_n is seen to be

$$s_n = \lambda_1 \lambda_2 \cdots \lambda_n = \det \mathbf{A} \quad (5.28)$$

Thus, $\text{tr } \mathbf{A}$ and $\det \mathbf{A}$ are both invariant to similarity transformations.

Let us now restrict ourselves to the case in which $\mathbf{A}$ is a self-adjoint matrix. Then the discussion becomes similar to that for the eigenvalue problems of waveguides and cavities. Let $\mathbf{x}_i$ and λ_i be an eigenvector and the corresponding eigenvalue, respectively. Then, we have by definition

$$\mathbf{A} \mathbf{x}_i - \lambda_i \mathbf{x}_i = 0 \quad (5.29)$$

Multiplying by $\mathbf{x}_i^+$ from the left, this becomes

$$\mathbf{x}_i^+ \mathbf{A} \mathbf{x}_i - \lambda_i \mathbf{x}_i^+ \mathbf{x}_i = 0 \quad (5.30)$$

Next, taking the adjoint matrix of (5.29) and then multiplying by $\mathbf{x}_i$ from the right, we obtain

$$\mathbf{x}_i^+ \mathbf{A}^+ \mathbf{x}_i - \lambda_i^* \mathbf{x}_i^+ \mathbf{x}_i = 0 \quad (5.31)$$

Since $\mathbf{A}$ is self-adjoint, the first terms in (5.30) and (5.31) are equal to each other. Subtracting (5.31) from (5.30), we have

$$(\lambda_i^* - \lambda_i) \mathbf{x}_i^+ \mathbf{x}_i = 0 \quad (5.32)$$

which shows that $\lambda_i = \lambda_i^*$, since $\mathbf{x}_i^+ \mathbf{x}_i \neq 0$. In other words, the eigenvalues of a self-adjoint matrix are real.

Let $\mathbf{x}_i$ and $\mathbf{x}_j$ be two eigenvectors with different eigenvalues. Multiplying (5.29) by $\mathbf{x}_j^+$ from the left, we have

$$\mathbf{x}_j^+ \mathbf{A} \mathbf{x}_i - \lambda_i \mathbf{x}_j^+ \mathbf{x}_i = 0 \quad (5.33)$$

After changing the subscript i in (5.29) to j and taking the adjoint matrix, if we multiply by $\mathbf{x}_i$ from the right, we obtain

$$\mathbf{x}_j^+ \mathbf{A}^+ \mathbf{x}_i - \lambda_j^* \mathbf{x}_j^+ \mathbf{x}_i = 0 \quad (5.34)$$

The first terms in (5.33) and (5.34) are equal to each other. Subtracting (5.33) from (5.34), we have

$$(\lambda_i - \lambda_j^*) \mathbf{x}_j^+ \mathbf{x}_i = 0 \quad (5.35)$$

Since λ_j is real, λ_j^* can be replaced by λ_j which is different from λ_i by hypothesis. Thus, it is shown that

$$\mathbf{x}_j^+ \mathbf{x}_i = 0 \quad (5.36)$$

This is called the orthogonality relation between $\mathbf{x}_i$ and $\mathbf{x}_j$. The orthogonal relation between two vectors in real three-dimensional space can be expressed in the same form, as is easily seen from (2.3) and (2.6) by setting $\theta = 90^\circ$.

Let us assume that all the eigenvectors are multiplied by appropriate constants so as to satisfy

$$\mathbf{x}_i^+ \mathbf{x}_i = \mathbf{I} \quad (5.37)$$

This is the normalization condition corresponding to (3.37) or (3.74). The normalization process, of course, does not change $\mathbf{x}_i$ from being an eigenvector.

A matrix of order 1×1 is different from an ordinary number, but treating this matrix as if it were a number whose value is equal to that of the only component, Eq. (5.30) gives

$$\lambda_i = \frac{\mathbf{x}_i^+ \mathbf{A} \mathbf{x}_i}{\mathbf{x}_i^+ \mathbf{x}_i} \quad (5.38)$$

This suggests that

$$\lambda(\mathbf{x}) = \frac{\mathbf{x}^+ \mathbf{A} \mathbf{x}}{\mathbf{x}^+ \mathbf{x}} \quad (5.39)$$

may be a variational expression for the eigenvalues. Since A is self-adjoint, we have

$$\lambda^*(\mathbf{x}) = \frac{(\mathbf{x}^+ \mathbf{A} \mathbf{x})^+}{(\mathbf{x}^+ \mathbf{x})^+} = \frac{\mathbf{x}^+ \mathbf{A}^+ \mathbf{x}}{\mathbf{x}^+ \mathbf{x}} = \frac{\mathbf{x}^+ \mathbf{A} \mathbf{x}}{\mathbf{x}^+ \mathbf{x}} = \lambda(\mathbf{x})$$

which shows that $\lambda(\mathbf{x})$ is real, regardless of the value of nonzero $\mathbf{x}$. Suppose that λ becomes $\lambda + \delta\lambda$ when $\mathbf{x}$ becomes $\mathbf{x} + \delta\mathbf{x}$. Multiplying (5.39) by the denominator on the right-hand side and taking the variation, we obtain

$$\lambda(\delta\mathbf{x}^+ \mathbf{x} + \mathbf{x}^+ \delta\mathbf{x}) + \delta\lambda \mathbf{x}^+ \mathbf{x} = \delta\mathbf{x}^+ \mathbf{A} \mathbf{x} + \mathbf{x}^+ \mathbf{A} \delta\mathbf{x}$$

A little manipulation shows that this is equivalent to

$$\delta\lambda \mathbf{x}^+ \mathbf{x} = \delta\mathbf{x}^+ (\mathbf{A} \mathbf{x} - \lambda \mathbf{x}) + (\mathbf{A} \mathbf{x} - \lambda \mathbf{x})^+ \delta\mathbf{x} \quad (5.40)$$

where we have used $A = A^+$ and $\lambda = \lambda^*$. If $\mathbf{x}$ is an eigenvector, from (5.40), the first order variation $\delta\lambda$ of λ becomes zero. Conversely, if $\delta\lambda = 0$ for all possible first order variations $\delta\mathbf{x}$ from $\mathbf{x}$, (5.40) shows that the real part of $\delta\mathbf{x}^+ (A\mathbf{x} - \lambda\mathbf{x})$ must be zero, regardless of the direction of $\delta\mathbf{x}^+$. However, if $A\mathbf{x} - \lambda\mathbf{x}$ is not a $\mathbf{0}$ matrix, by choosing a proper $\delta\mathbf{x}^+$, the real part of $\delta\mathbf{x}^+ (A\mathbf{x} - \lambda\mathbf{x})$ can be made nonzero leading to a contradiction. Therefore, if $\delta\lambda = 0$ for all possible variations $\delta\mathbf{x}$ from $\mathbf{x}$, we can conclude that $\mathbf{x}$ is an eigenvector and that $\lambda(\mathbf{x})$ gives the corresponding eigenvalue. This completes the proof that (5.39) is, indeed, a variational expression for the eigenvalues.

Once the variational expression for λ has been obtained, all the eigenvectors can be found sequentially, at least conceptionally. For example, we first obtain a vector $\mathbf{x}_1$ which minimizes $\lambda(\mathbf{x})$. Then, we obtain $\mathbf{x}_2$ which minimizes $\lambda(\mathbf{x})$ under the restriction of being orthogonal to $\mathbf{x}_1$. Since $\delta\lambda = 0$ for all $\delta\mathbf{x}$, $\mathbf{x}_1$ is obviously an eigenvector; $\mathbf{x}_2$ is also an eigenvector which can be shown as follows. An arbitrary vector $\mathbf{x}$ can be written in the form

$$\mathbf{x} = (\mathbf{x}_1^+ \mathbf{x}) \mathbf{x}_1 + \{\mathbf{x} - (\mathbf{x}_1^+ \mathbf{x}) \mathbf{x}_1\} \quad (5.41)$$

This shows that $\mathbf{x}$ can be expressed as the sum of two vectors, one parallel to $\mathbf{x}_1$, and the other orthogonal to $\mathbf{x}_1$. In that part of $\delta\mathbf{x}$ which is orthogonal to $\mathbf{x}_1$, $\delta\lambda$ is equal to zero, as can be seen from the method of obtaining $\mathbf{x}_2$. In that part of $\delta\mathbf{x}$ which is parallel to $\mathbf{x}_1$, $\delta\mathbf{x}_1$,

$$\delta\lambda \mathbf{x}_2^+ \mathbf{x}_2 = \delta\mathbf{x}_1^+ (A\mathbf{x}_2 - \lambda\mathbf{x}_2) + (A\mathbf{x}_2 - \lambda\mathbf{x}_2)^+ \delta\mathbf{x}_1$$

Since $\delta\mathbf{x}_1$ is now orthogonal to $\mathbf{x}_2$, $\delta\mathbf{x}_1^+ \lambda\mathbf{x}_2$ and $\lambda^* \mathbf{x}_2^+ \delta\mathbf{x}_1$ are both equal to zero. Furthermore, we have

$$\delta\mathbf{x}_1^+ A\mathbf{x}_2 = (A \delta\mathbf{x}_1)^+ \mathbf{x}_2 = \lambda_1 \delta\mathbf{x}_1^+ \mathbf{x}_2 = 0$$

Similarly, $(A\mathbf{x}_2)^+ \delta\mathbf{x}_1$ is equal to zero. As a result, $\delta\lambda = 0$ for $\delta\mathbf{x}_1$. We can conclude, therefore, that $\delta\lambda$ is equal to zero, regardless of the direction of $\delta\mathbf{x}$, and hence $\mathbf{x}_2$ is an eigenvector. Similarly, under the restriction of being orthogonal to $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{i-1}$, if we obtain a vector $\mathbf{x}_i$ which minimizes $\lambda(\mathbf{x})$, then $\mathbf{x}_i$ becomes the i th eigenvector. When i becomes n , the number of the dimensions of the space under consideration, no further eigenvectors can be obtained, for in the n -dimensional space, there is no nonzero vector which is orthogonal to n orthogonal vectors. Therefore, the above procedure gives exactly n eigenvectors.

As we stated in connection with (5.37), the eigenvectors are all assumed to be normalized. An arbitrary vector $\mathbf{x}$ in the n -dimensional space can

therefore be expressed in the form

$$\mathbf{x} = \sum_{i=1}^n (\mathbf{x}_i^+ \mathbf{x}) \mathbf{x}_i \quad (5.42)$$

which can be checked by multiplying $\mathbf{x}_i^+$ ($i=1, 2, \dots, n$) from the left. Multiplying (5.42) by $\mathbf{A}$ from the left, we obtain

$$\mathbf{A}\mathbf{x} = \sum_{i=1}^n (\mathbf{x}_i^+ \mathbf{x}) \lambda_i \mathbf{x}_i \quad (5.43)$$

where $\mathbf{A}\mathbf{x}_i = \lambda_i \mathbf{x}_i$ is used. From this, $\mathbf{A}$ can be written in the form

$$\mathbf{A} \cdot = \sum_{i=1}^n \lambda_i \mathbf{x}_i \mathbf{x}_i^+ \quad (5.44)$$

This is called the spectral representation of $\mathbf{A}$ where the dot signifies that $\mathbf{A}$ is an operator. When none of the eigenvalues become zero, $\mathbf{A}^{-1}$ exists and is given by

$$\mathbf{A}^{-1} \cdot = \sum_{i=1}^n \lambda_i^{-1} \mathbf{x}_i \mathbf{x}_i^+$$

The validity of this expression can be checked by substituting into (5.8).

Let $\mathbf{X}$ be a matrix constructed by the eigenvectors $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ of a self-adjoint matrix $\mathbf{A}$ in the form

$$\mathbf{X} = [\mathbf{x}_1 \quad \mathbf{x}_2 \quad \dots \quad \mathbf{x}_n]$$

Note that $\mathbf{X}$ is a square matrix. Because of the orthogonality condition between the eigenvectors, we have

$$\mathbf{X}^+ \mathbf{X} = \begin{bmatrix} \mathbf{x}_1^+ \\ \mathbf{x}_2^+ \\ \vdots \\ \mathbf{x}_n^+ \end{bmatrix} [\mathbf{x}_1 \quad \mathbf{x}_2 \quad \dots \quad \mathbf{x}_n] = \begin{bmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & & \vdots \\ 0 & 0 & \dots & 1 \end{bmatrix} = \mathbf{I} \quad (5.45)$$

Taking the determinant of (5.45), the left-hand side becomes $\det \mathbf{X}^+ \det \mathbf{X}$, and the right-hand side becomes unity. From this we see that $\det \mathbf{X}$ cannot be equal to zero, and hence $\mathbf{X}^{-1}$ exists. Multiplying (5.45) by $\mathbf{X}^{-1}$ from the right, we obtain

$$\mathbf{X}^+ = \mathbf{X}^{-1} \quad (5.46)$$

which shows that $\mathbf{X}$ is a unitary matrix. Multiplying $\mathbf{A}$ by $\mathbf{X}^+$ and $\mathbf{X}$ from

the left and right, respectively, we have

$$\begin{aligned} \mathbf{X}^+ \mathbf{A} \mathbf{X} &= \mathbf{X}^+ [\lambda_1 \mathbf{x}_1 \quad \lambda_2 \mathbf{x}_2 \quad \dots \quad \lambda_n \mathbf{x}_n] \\ &= \begin{bmatrix} \lambda_1 & 0 & \dots & 0 \\ 0 & \lambda_2 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & \lambda_n \end{bmatrix} = \text{diag}[\lambda_1 \quad \lambda_2 \quad \dots \quad \lambda_n] \end{aligned} \quad (5.47)$$

where $\mathbf{A} \mathbf{x}_i = \lambda_i \mathbf{x}_i$ is used. The symbol on the right-hand side indicates the diagonal matrix with diagonal components $\lambda_1, \lambda_2, \dots, \lambda_n$. Using (5.46), the left-hand side of (5.47) is seen to be the similarity transformation of $\mathbf{A}$ by $\mathbf{X}^{-1}$. Thus, all the eigenvalues of $\mathbf{A}$ can be obtained from the similarity transformation by $\mathbf{X}^{-1}$.

Now suppose that $\mathbf{x}^+ \mathbf{B} \mathbf{x}$ vanishes regardless of the value of $\mathbf{x}$ (i.e., regardless of the values of the components of $\mathbf{x}$), then we have

$$\sum_{ij} x_i^* B_{ij} x_j = 0$$

If we assume that the components of $\mathbf{x}$ are all zero, except for x_i and x_j , the above equation reduces to

$$x_i^* B_{ij} x_j + x_j^* B_{ji} x_i = 0$$

Setting $x_i = x_j = 1$, this becomes

$$B_{ij} + B_{ji} = 0$$

Similarly, setting $x_i = j$ and $x_j = 1$, we obtain

$$-B_{ij} + B_{ji} = 0$$

These two equations show that $B_{ij} = B_{ji} = 0$, and since i and j are arbitrary, we can conclude that if $\mathbf{x}^+ \mathbf{B} \mathbf{x}$ vanishes, regardless of the value of $\mathbf{x}$, $\mathbf{B}$ must be a $\mathbf{0}$ matrix.

Next, suppose that $\mathbf{x}^+ \mathbf{A} \mathbf{x}$ is real regardless of the value of $\mathbf{x}$. Since $(\mathbf{x}^+ \mathbf{A} \mathbf{x})^+$ has the same value as $\mathbf{x}^+ \mathbf{A} \mathbf{x}$, we have $\mathbf{x}^+ \mathbf{A} \mathbf{x} = \mathbf{x}^+ \mathbf{A}^+ \mathbf{x}$, or equivalently

$$\mathbf{x}^+ (\mathbf{A} - \mathbf{A}^+) \mathbf{x} = 0$$

This equation holds regardless of the value of $\mathbf{x}$, hence $\mathbf{A} - \mathbf{A}^+$ must be a $\mathbf{0}$ matrix; in other words, $\mathbf{A}$ is self-adjoint ($\mathbf{A} = \mathbf{A}^+$). Conversely, if $\mathbf{A}$ is self-adjoint, $\mathbf{x}^+ \mathbf{A} \mathbf{x} = (\mathbf{x}^+ \mathbf{A} \mathbf{x})^+$, and $\mathbf{x}^+ \mathbf{A} \mathbf{x}$ becomes real regardless of the

value of $\mathbf{x}$. Thus, $\mathbf{A}$ is a self-adjoint matrix if, and only if, $\mathbf{x}^+ \mathbf{A} \mathbf{x}$ is real, regardless of the value of $\mathbf{x}$.

If $\mathbf{x}^+ \mathbf{A} \mathbf{x}$ is always positive for nonzero $\mathbf{x}$, $\mathbf{A}$ is said to be positive-definite. If $\mathbf{A}$ is positive-definite, it is, of course, self-adjoint for positive numbers are real. It is also obvious from (5.39) that, if $\mathbf{A}$ is positive-definite, the eigenvalues are all positive. Conversely, if the eigenvalues are all positive, $\mathbf{A}$ is positive-definite. This can be seen as follows from (5.42) and (5.43)

$$\mathbf{x}^+ \mathbf{A} \mathbf{x} = \sum_i (\mathbf{x}^+ \mathbf{x}_i) \lambda_i (\mathbf{x}_i^+ \mathbf{x}) = \sum_i \lambda_i |\mathbf{x}_i^+ \mathbf{x}|^2$$

which is always positive for nonzero $\mathbf{x}$. Thus, $\mathbf{A}$ is positive-definite if, and only if, the eigenvalues are all positive.

Suppose that $\mathbf{A}$ is positive-definite, then a square matrix $\mathbf{T}$ can be defined by

$$\mathbf{T} = [\lambda_1^{-1/2} \mathbf{x}_1 \quad \lambda_2^{-1/2} \mathbf{x}_2 \quad \dots \quad \lambda_n^{-1/2} \mathbf{x}_n]$$

where $\lambda_i \neq 0$ ($i=1, 2, \dots, n$) is used. It can easily be seen that $\mathbf{T}$ is non-singular since $\mathbf{T}^+ \mathbf{T} = \text{diag}[\lambda_1^{-1} \lambda_2^{-1} \dots \lambda_n^{-1}]$, and $\det \mathbf{T} \neq 0$. Furthermore, using (5.47), we have

$$\mathbf{T}^+ \mathbf{A} \mathbf{T} = \mathbf{I}$$

Conversely, suppose that $\mathbf{T}^+ \mathbf{A} \mathbf{T}$ forms a unit matrix for a particular $\mathbf{T}$, then

$$\mathbf{y}^+ \mathbf{T}^+ \mathbf{A} \mathbf{T} \mathbf{y} = \mathbf{y}^+ \mathbf{y}$$

is always positive for nonzero $\mathbf{y}$. Setting $\mathbf{x} = \mathbf{T} \mathbf{y}$ and noting that $\mathbf{x}$ can take any nonzero value, since $\mathbf{y}$ is an arbitrary nonzero vector and $\mathbf{T}$ is non-singular, $\mathbf{x}^+ \mathbf{A} \mathbf{x}$ is then seen to be positive for any nonzero $\mathbf{x}$. Thus, $\mathbf{A}$ is positive-definite if, and only if, there exists a nonsingular matrix $\mathbf{T}$ which makes $\mathbf{T}^+ \mathbf{A} \mathbf{T}$ a unit matrix.

Another necessary and sufficient condition for a self-adjoint matrix $\mathbf{A}$ to be positive-definite is that the components of $\mathbf{A}$ satisfy

$$a_{11} > 0, \quad \begin{vmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{vmatrix} > 0, \dots, \quad \begin{vmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{vmatrix} > 0 \quad (5.48)$$

The conditions given in (5.48) do not require the knowledge of the eigenvalues, and hence they provide a practical method for checking the positive-definiteness of a self-adjoint matrix. Although the proof is lengthy, it will be given below.

When $\mathbf{A}$ is a square matrix of order 1, the necessary and sufficient condi-

tion for $\mathbf{A}$ to be positive-definite is obviously given by $a_{11} > 0$. It will be appropriate, therefore, to attempt a proof for the general case by mathematical induction. First, let us prove the sufficiency. Let $\mathbf{A}_{n-1}$ be a square matrix of order $n-1$ with the first $n-1$ rows and $n-1$ columns of $\mathbf{A}$. When the first $n-1$ conditions in (5.48) are satisfied, assume that $\mathbf{A}_{n-1}$ is positive-definite; i.e., there exists a nonsingular matrix $\mathbf{T}_{n-1}$ satisfying

$$\mathbf{T}_{n-1}^+ \mathbf{A}_{n-1} \mathbf{T}_{n-1} = \mathbf{I} \quad (5.49)$$

With this assumption, if $\mathbf{A}$ can be proved to be positive-definite using the last condition in (5.48), then (5.48) is sufficient. Let us write $\mathbf{x}$ and $\mathbf{A}$ in the forms

$$\mathbf{x} = \begin{bmatrix} \mathbf{x}_{n-1} \\ x_n \end{bmatrix}, \quad \mathbf{A} \equiv \mathbf{A}_n = \begin{bmatrix} \mathbf{A}_{n-1} & \mathbf{a}_{n-1} \\ \mathbf{a}_{n-1}^+ & a_{nn} \end{bmatrix}$$

then we have

$$\begin{aligned} \mathbf{x}^+ \mathbf{A} \mathbf{x} &= [\mathbf{x}_{n-1}^+ \quad x_n^*] \begin{bmatrix} \mathbf{A}_{n-1} & \mathbf{a}_{n-1} \\ \mathbf{a}_{n-1}^+ & a_{nn} \end{bmatrix} \begin{bmatrix} \mathbf{x}_{n-1} \\ x_n \end{bmatrix} \\ &= \mathbf{x}_{n-1}^+ \mathbf{A}_{n-1} \mathbf{x}_{n-1} + x_n \mathbf{x}_{n-1}^+ \mathbf{a}_{n-1} + x_n^* \mathbf{a}_{n-1}^+ \mathbf{x}_{n-1} + a_{nn} x_n^* x_n \end{aligned} \quad (5.50)$$

Let $\mathbf{y}$ be a vector defined by means of the relation

$$\mathbf{x} = \begin{bmatrix} \mathbf{x}_{n-1} \\ x_n \end{bmatrix} = \begin{bmatrix} \mathbf{T}_{n-1} & -\mathbf{T}_{n-1} \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1} \\ 0 & 1 \end{bmatrix} \begin{bmatrix} \mathbf{y}_{n-1} \\ y_n \end{bmatrix} \equiv \mathbf{L}_n \mathbf{y} = \mathbf{L} \mathbf{y} \quad (5.51)$$

Since $\det \mathbf{L} = \det \mathbf{T}_{n-1}$ and $\mathbf{T}_{n-1}$ is nonsingular, $\mathbf{L}$ is also nonsingular. With this $\mathbf{y}$, (5.50) can be rewritten in the form

$$\begin{aligned} \mathbf{y}^+ \mathbf{L}^+ \mathbf{A} \mathbf{L} \mathbf{y} &= (\mathbf{y}_{n-1} - \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1} y_n)^+ \mathbf{T}_{n-1}^+ \mathbf{A}_{n-1} \mathbf{T}_{n-1} (\mathbf{y}_{n-1} - \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1} y_n) \\ &\quad + y_n (\mathbf{y}_{n-1} - \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1} y_n)^+ \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1} \\ &\quad + y_n^* \mathbf{a}_{n-1}^+ \mathbf{T}_{n-1} (\mathbf{y}_{n-1} - \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1} y_n) + a_{nn} y_n^* y_n \end{aligned}$$

The right-hand side can be simplified using (5.49) which results in

$$\mathbf{y}^+ \mathbf{L}^+ \mathbf{A} \mathbf{L} \mathbf{y} = y_{n-1}^+ y_{n-1} + (a_{nn} - \mathbf{a}_{n-1}^+ \mathbf{T}_{n-1} \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1}) y_n^* y_n \quad (5.52)$$

This shows that $\mathbf{L}^+ \mathbf{A} \mathbf{L}$ is a diagonal matrix given by

$$\mathbf{L}^+ \mathbf{A} \mathbf{L} = \text{diag}[1 \quad 1 \dots 1 \quad a_{nn} - \mathbf{a}_{n-1}^+ \mathbf{T}_{n-1} \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1}] \quad (5.53)$$

where the first $n-1$ diagonal components are all unity. Taking the determinant of (5.53), we have

$$(\det \mathbf{L}^+) (\det \mathbf{A}) (\det \mathbf{L}) = |\det \mathbf{L}|^2 \det \mathbf{A} = a_{nn} - \mathbf{a}_{n-1}^+ \mathbf{T}_{n-1} \mathbf{T}_{n-1}^+ \mathbf{a}_{n-1} \quad (5.54)$$

From the last condition in (5.48), $\det \mathbf{A}$ is positive; hence, the right-hand side of (5.54) must be positive. Again, introducing a new vector $\mathbf{z}$ defined by

$$\mathbf{z} = \begin{bmatrix} \mathbf{I} & 0 \\ 0 & (a_{mm} - \mathbf{a}_{m-1}^+ \mathbf{T}_{m-1} \mathbf{T}_{m-1}^+ \mathbf{a}_{m-1})^{1/2} \end{bmatrix} \begin{bmatrix} \mathbf{y}_{m-1} \\ y_m \end{bmatrix} = \mathbf{M}^{-1} \mathbf{y} \quad (5.55)$$

where $\mathbf{M}^{-1}$ is clearly nonsingular, (5.52) becomes

$$\mathbf{z}^+ \mathbf{M}^+ \mathbf{L}^+ \mathbf{A} \mathbf{L} \mathbf{M} \mathbf{z} = \mathbf{z}^+ \mathbf{z} \quad (5.56)$$

Noting that $\mathbf{x}$ is an arbitrary nonzero vector, which means that $\mathbf{y}$ and $\mathbf{z}$ are both arbitrary, (5.56) indicates that

$$\mathbf{M}^+ \mathbf{L}^+ \mathbf{A} \mathbf{L} \mathbf{M} = \mathbf{I} \quad (5.57)$$

If we set $\mathbf{L} \mathbf{M} = \mathbf{T}$, then since $\mathbf{T}$ is nonsingular, (5.57) shows that a nonsingular matrix $\mathbf{T}$ exists which makes $\mathbf{T}^+ \mathbf{A} \mathbf{T}$ a unit matrix. Therefore, $\mathbf{A}$ is positive-definite, and the proof of sufficiency is complete.

Let us next consider the necessity of (5.48). If $a_{11} \leq 0$, then $\mathbf{x}^+ \mathbf{A} \mathbf{x}$ cannot be positive for a particular $\mathbf{x}^+$ given by $[1 \ 0 \ 0 \dots 0]$. Therefore, $a_{11} > 0$ is a necessary condition for $\mathbf{A}$ to be positive-definite. To apply the principle of mathematical induction, let us assume that the first $m-1$ conditions in (5.48) hold, but the m th condition does not. Then, from the sufficiency of the first $m-1$ conditions, there exists a nonsingular matrix $\mathbf{T}_{m-1}$ satisfying

$$\mathbf{T}_{m-1}^+ \mathbf{A}_{m-1} \mathbf{T}_{m-1} = \mathbf{I} \quad (5.58)$$

Define a square matrix $\mathbf{L}_m$ similar to $\mathbf{L}_n$ in (5.51) using $\mathbf{T}_{m-1}$ and $\mathbf{a}_{m-1}$ in place of $\mathbf{T}_{n-1}$ and $\mathbf{a}_{n-1}$, respectively. Then, $\mathbf{L}_m$ is nonsingular, and we have

$$\mathbf{L}_m^+ \mathbf{A}_m \mathbf{L}_m = \text{diag}[1 \ 1 \ \dots \ 1 \ a_{mm} - \mathbf{a}_{m-1}^+ \mathbf{T}_{m-1} \mathbf{T}_{m-1}^+ \mathbf{a}_{m-1}] \quad (5.59)$$

If $\det \mathbf{A}_m \leq 0$, $a_{mm} - \mathbf{a}_{m-1}^+ \mathbf{T}_{m-1} \mathbf{T}_{m-1}^+ \mathbf{a}_{m-1}$ can be proved nonpositive (negative or zero) using an expression similar to (5.54). Let $\mathbf{x}_m = \mathbf{L}_m \mathbf{y}_m$ and $\mathbf{y}_m^+ = [0 \ 0 \dots 0 \ 1]$. Then, for $\mathbf{x}^+$ given by

$$\mathbf{x}^+ = [\mathbf{x}_m^+ \ 0 \ 0 \ \dots \ 0]$$

where the last $n-m$ components are all zero, we have

$$\mathbf{x}^+ \mathbf{A} \mathbf{x} = \mathbf{x}_m^+ \mathbf{A}_m \mathbf{x}_m = \mathbf{y}_m^+ \mathbf{L}_m^+ \mathbf{A}_m \mathbf{L}_m \mathbf{y}_m = a_{mm} - \mathbf{a}_{m-1}^+ \mathbf{T}_{m-1} \mathbf{T}_{m-1}^+ \mathbf{a}_{m-1}$$

which is nonpositive. Therefore, the m th condition is also necessary for the positive definiteness. Since m is arbitrary, all the conditions in (5.48) are necessary, and the proof is complete.

If the range of $\mathbf{x}^+ \mathbf{A} \mathbf{x}$ for arbitrary nonzero $\mathbf{x}$ includes zero and otherwise remains positive, $\mathbf{A}$ is said to be semipositive-definite. Matrix $\mathbf{A}$ is semipositive-definite if, and only if, at least one of the eigenvalues is zero and the remainder are not negative. From (5.28), $\det \mathbf{A}$ is equal to zero in this case, and $\mathbf{A}^{-1}$ does not exist.

Now, suppose that there are two self-adjoint matrices, $\mathbf{A}$ and $\mathbf{B}$, and $\mathbf{A}$ is positive-definite, then from the positive-definiteness of $\mathbf{A}$, there is a non-singular matrix $\mathbf{T}$ which satisfies $\mathbf{T}^+ \mathbf{A} \mathbf{T} = \mathbf{I}$. Since

$$(\mathbf{T}^+ \mathbf{B} \mathbf{T})^+ = \mathbf{T}^+ \mathbf{B}^+ \mathbf{T} = \mathbf{T}^+ \mathbf{B} \mathbf{T}$$

then $\mathbf{T}^+ \mathbf{B} \mathbf{T}$ is self-adjoint. Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ be the normalized eigenvectors of $\mathbf{T}^+ \mathbf{B} \mathbf{T}$, and let $\mathbf{X}$ be a square matrix constructed from them, i.e., $\mathbf{X} = [\mathbf{x}_1 \mathbf{x}_2 \dots \mathbf{x}_n]$. Then, $\mathbf{X}^+ \mathbf{T}^+ \mathbf{B} \mathbf{T} \mathbf{X}$ must be $\text{diag}[\lambda_1 \lambda_2 \dots \lambda_n]$ where $\lambda_i (i = 1, 2, \dots, n)$ is the eigenvalue corresponding to $\mathbf{x}_i$. On the other hand, we obviously have

$$\mathbf{X}^+ \mathbf{T}^+ \mathbf{A} \mathbf{T} \mathbf{X} = \mathbf{X}^+ \mathbf{I} \mathbf{X} = \mathbf{X}^+ \mathbf{X} = \mathbf{I}$$

therefore, if we define $\mathbf{H}$ by $\mathbf{H} = \mathbf{T} \mathbf{X}$, $\mathbf{H}$ is nonsingular and

$$\mathbf{H}^+ \mathbf{A} \mathbf{H} = \mathbf{I} \quad (5.60)$$

$$\mathbf{H}^+ \mathbf{B} \mathbf{H} = \text{diag}[\lambda_1 \quad \lambda_2 \quad \dots \quad \lambda_n] \quad (5.61)$$

The above discussion shows that if there are two self-adjoint matrices $\mathbf{A}$ and $\mathbf{B}$, and one of them, say $\mathbf{A}$, is positive-definite, then there is a non-singular matrix $\mathbf{H}$ which diagonalizes both $\mathbf{H}^+ \mathbf{A} \mathbf{H}$ and $\mathbf{H}^+ \mathbf{B} \mathbf{H}$, simultaneously; i.e., all of their components except those on the main diagonal are made zero. This fact will be used in the noise discussion of linear amplifiers in Chapter 7.

In the above discussion, $\mathbf{H}^+$ is not necessarily equal to $\mathbf{H}^{-1}$. From (5.60), we have

$$\mathbf{H}^+ = (\mathbf{A} \mathbf{H})^{-1} = \mathbf{H}^{-1} \mathbf{A}^{-1}$$

Substituting this into (5.61), we obtain

$$\mathbf{H}^{-1} \mathbf{A}^{-1} \mathbf{B} \mathbf{H} = \text{diag}[\lambda_1 \quad \lambda_2 \quad \dots \quad \lambda_n]$$

The similarity transformation of $\mathbf{A}^{-1} \mathbf{B}$ by $\mathbf{H}^{-1}$ gives the eigenvalues $\lambda_1, \lambda_2, \dots, \lambda_n$. However, since similarity transformations do not change the eigenvalues of a matrix, as explained in connection with (5.24) and (5.25), $\lambda_1, \lambda_2, \dots, \lambda_n$ must also be eigenvalues of $\mathbf{A}^{-1} \mathbf{B}$. Consequently, $\lambda_1, \lambda_2, \dots, \lambda_n$ in (5.61) can be obtained by calculating the eigenvalues of $\mathbf{A}^{-1} \mathbf{B}$. Both $\mathbf{A}$

and $\mathbf{B}$ are self-adjoint, but $\mathbf{A}^{-1}\mathbf{B}$ is not necessarily so. Nevertheless, its eigenvalues are all real because $\lambda_1, \lambda_2, \dots, \lambda_n$ are the eigenvalues of a self-adjoint matrix $\mathbf{T}^+\mathbf{B}\mathbf{T}$.

Finally, let us consider a necessary and sufficient condition for two different self-adjoint matrices $\mathbf{A}$ and $\mathbf{B}$ to be simultaneously diagonalized by a unitary matrix through the similarity transformation. Although the discussion is instructive and worth presenting, we shall not utilize the result in this book, and one may omit the remainder of this section on a first reading and proceed to Section 5.2.

Suppose that the similarity transformations of two different self-adjoint matrices $\mathbf{A}$ and $\mathbf{B}$ by a unitary matrix give diagonal matrices $\mathbf{A}'$ and $\mathbf{B}'$, respectively. The order of the product of two diagonal matrices can be interchanged without changing the result; i.e.,

$$\mathbf{A}'\mathbf{B}' = \mathbf{B}'\mathbf{A}'$$

The forms of matrix equations are invariant to similarity transformations, and hence we have

$$\mathbf{AB} = \mathbf{BA} \quad (5.62)$$

The above discussion shows that (5.62) is a necessary condition for $\mathbf{A}$ and $\mathbf{B}$ to be simultaneously diagonalized by a unitary matrix through the similarity transformation. The proof that (5.62) is also sufficient will be given below. We first assume that (5.62) holds, and prove that there is a unitary matrix which diagonalizes both $\mathbf{A}$ and $\mathbf{B}$ through the similarity transformation. Let $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n$ be the normalized eigenvectors of $\mathbf{A}$, and let $\mathbf{X}$ be a matrix constructed from them; i.e.,

$$\mathbf{X} = [\mathbf{x}_1 \quad \mathbf{x}_2 \quad \dots \quad \mathbf{x}_n]$$

Then, we have

$$\mathbf{X}^{-1}\mathbf{AX} = \mathbf{X}^+\mathbf{AX} = \text{diag} [\lambda_1 \quad \lambda_2 \quad \dots \quad \lambda_n] \quad (5.63)$$

where λ_i is the eigenvalue corresponding to $\mathbf{x}_i$. For convenience, we assume that $\lambda_1 \leq \lambda_2 \leq \lambda_3 \leq \dots \leq \lambda_n$. In order to calculate $\mathbf{X}^{-1}\mathbf{BX}$, let us consider $\mathbf{x}_i^+\mathbf{ABx}_j$. From (5.62), $\mathbf{x}_i^+\mathbf{ABx}_j$ is equal to $\mathbf{x}_i^+\mathbf{BAX}_j$. From the fact that both $\mathbf{A}$ and $\mathbf{B}$ are self-adjoint, we have

$$\mathbf{x}_i^+\mathbf{ABx}_j = (\mathbf{Ax}_i)^+ \mathbf{Bx}_j = \lambda_i \mathbf{x}_i^+ \mathbf{Bx}_j = \mathbf{x}_i^+ \mathbf{BAX}_j = \mathbf{x}_i^+ \mathbf{B}\lambda_j \mathbf{x}_j = \lambda_j \mathbf{x}_i^+ \mathbf{Bx}_j$$

It follows from this that, if $\lambda_i \neq \lambda_j$,

$$\mathbf{x}_i^+ \mathbf{Bx}_j = 0$$

In other words, $\mathbf{X}^+ \mathbf{B} \mathbf{X}$ must have the form

$$\mathbf{X}^+ \mathbf{B} \mathbf{X} = \mathbf{X}^{-1} \mathbf{B} \mathbf{X} = \begin{bmatrix} \mathbf{B}_1 & 0 & 0 & \dots \\ 0 & \mathbf{B}_2 & 0 & \dots \\ 0 & 0 & \mathbf{B}_3 & \dots \\ \vdots & \vdots & & \ddots \end{bmatrix}$$

where $\mathbf{B}_l$ ($l = 1, 2, \dots$) is a square matrix with the number of columns (and rows) equal to the degree of the degeneracy of the l th smallest eigenvalue of $\mathbf{A}$, where the l th smallest eigenvalue is such that there are $l - 1$ discrete eigenvalues smaller than it. The largest subscript is equal to the number of different eigenvalues of $\mathbf{A}$. Since $(\mathbf{X}^+ \mathbf{B} \mathbf{X})^+ = \mathbf{X}^+ \mathbf{B}^+ \mathbf{X} = \mathbf{X}^+ \mathbf{B} \mathbf{X}$, we have

$$\mathbf{B}_l^+ = \mathbf{B}_l$$

which shows that the $\mathbf{B}_l$'s are also self-adjoint. Let $\mathbf{Y}_l$ be a unitary matrix constructed from the normalized eigenvectors of $\mathbf{B}_l$ in an usual manner. Then the similarity transformation of $\mathbf{B}_l$ by $\mathbf{Y}_l^+$ gives a diagonal matrix. Let $\mathbf{Y}$ be a square matrix defined by

$$\mathbf{Y} = \begin{bmatrix} \mathbf{Y}_1 & 0 & 0 & \dots \\ 0 & \mathbf{Y}_2 & 0 & \dots \\ 0 & 0 & \mathbf{Y}_3 & \dots \\ \vdots & \vdots & & \ddots \end{bmatrix}$$

Then $\mathbf{Y}$ itself becomes a unitary matrix. From the method of constructing $\mathbf{Y}$, it is obvious that $\mathbf{Y}^+ \mathbf{X}^+ \mathbf{B} \mathbf{X} \mathbf{Y}$ is a diagonal matrix and that $\mathbf{Y}^+ \mathbf{X}^+ \mathbf{A} \mathbf{X} \mathbf{Y}$ becomes $\text{diag}[\lambda_1 \lambda_2 \dots \lambda_n]$ from (5.63). Noting that $\mathbf{X} \mathbf{Y}$ is a product of two unitary matrices, it must be a unitary matrix, and therefore we see that $(\mathbf{X} \mathbf{Y})^{-1} = \mathbf{Y}^+ \mathbf{X}^+$ is a unitary matrix. Thus, both $\mathbf{A}$ and $\mathbf{B}$ are shown to be diagonalized through the similarity transformations by the unitary matrix $\mathbf{Y}^+ \mathbf{X}^+$. This completes the proof.

5.2 Reciprocal Conditions

Suppose that several waveguides are connected together to form a waveguide junction as schematically shown in Fig. 5.1. At a reference plane far from the discontinuities in each waveguide, all the higher modes with frequencies below cutoff can be neglected, and only a finite number of propagating modes have to be considered. At such a reference plane, the voltage and current can be defined for each propagating mode as we explained in Chapter 3. Since Maxwell's equations indicate a linear relation between the

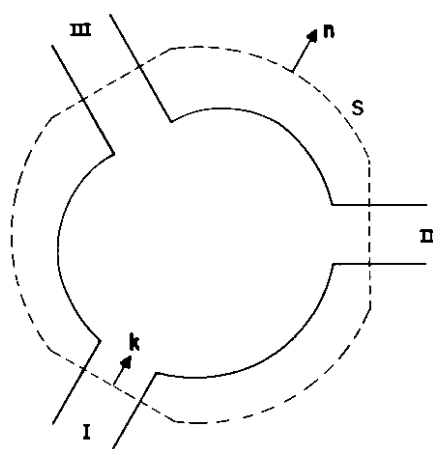


Fig. 5.1. Waveguide junction.

electric and magnetic fields, there should be a linear relation between the voltages and currents which represent the fields. Let us consider the case in which all currents are given independently, and the voltages are uniquely determined by them. In such a case, we have

$$\begin{aligned}
 V_1 &= Z_{11}I_1 + Z_{12}I_2 + \cdots + Z_{1n}I_n \\
 V_2 &= Z_{21}I_1 + Z_{22}I_2 + \cdots + Z_{2n}I_n \\
 &\vdots \\
 V_n &= Z_{n1}I_1 + Z_{n2}I_2 + \cdots + Z_{nn}I_n
 \end{aligned} \tag{5.64}$$

where V_i and I_i ($i = 1, 2, \dots, n$) are the voltage and current for a propagating mode at the reference plane in a waveguide shown in Fig. 5.1. When two or more propagating modes exist in some of the waveguides, n becomes larger than the number of actual openings of the junction. Since we are interested in the behavior of the waveguide junction as viewed from outside the reference planes, from (5.64) the junction as a whole can be considered as an n -port network. We can now use the concept of matrices explained in the previous section and define $\mathbf{v}$, $\mathbf{i}$, and $\mathbf{Z}$ by

$$\mathbf{v} = \begin{bmatrix} V_1 \\ V_2 \\ \vdots \\ V_n \end{bmatrix}, \quad \mathbf{i} = \begin{bmatrix} I_1 \\ I_2 \\ \vdots \\ I_n \end{bmatrix}, \quad \mathbf{Z} = \begin{bmatrix} Z_{11} & Z_{12} & \cdots & Z_{1n} \\ Z_{21} & Z_{22} & \cdots & Z_{2n} \\ \vdots & \vdots & & \vdots \\ Z_{n1} & Z_{n2} & \cdots & Z_{nn} \end{bmatrix} \tag{5.65}$$

respectively. Then (5.64) can be written in a particularly simple form,

$$\mathbf{v} = \mathbf{Z}\mathbf{i} \quad (5.66)$$

Similarly, if the voltages can be given independently of each other, and the currents are uniquely determined by them, we have

$$\mathbf{i} = \mathbf{Y}\mathbf{v} \quad (5.67)$$

If $\mathbf{Z}$ is nonsingular, we have from (5.66) and (5.67)

$$\mathbf{Z}^{-1} = \mathbf{Y} \quad (5.68)$$

In order to study relations between components of the matrix $\mathbf{Z}$, let $\mathbf{E}^{(1)}$ and $\mathbf{H}^{(1)}$ be the electric and magnetic fields inside the junction corresponding to a set of currents given at the ports. Similarly, let $\mathbf{E}^{(2)}$ and $\mathbf{H}^{(2)}$ be the fields corresponding to another set of currents. Using Maxwell's equations

$$\nabla \times \mathbf{H} = (\sigma + j\omega\epsilon) \mathbf{E}, \quad \nabla \times \mathbf{E} = -j\omega\mu\mathbf{H}$$

we have

$$\begin{aligned} \nabla \cdot (\mathbf{E}^{(1)} \times \mathbf{H}^{(2)} - \mathbf{E}^{(2)} \times \mathbf{H}^{(1)}) \\ = \mathbf{H}^{(2)} \cdot \nabla \times \mathbf{E}^{(1)} - \mathbf{E}^{(1)} \cdot \nabla \times \mathbf{H}^{(2)} - \mathbf{H}^{(1)} \cdot \nabla \times \mathbf{E}^{(2)} + \mathbf{E}^{(2)} \cdot \nabla \times \mathbf{H}^{(1)} \\ = -j\omega\mathbf{H}^{(2)} \cdot \mu\mathbf{H}^{(1)} - \mathbf{E}^{(1)} \cdot (\sigma + j\omega\epsilon)\mathbf{E}^{(2)} \\ + j\omega\mathbf{H}^{(1)} \cdot \mu\mathbf{H}^{(2)} + \mathbf{E}^{(2)} \cdot (\sigma + j\omega\epsilon)\mathbf{E}^{(1)} \end{aligned} \quad (5.69)$$

Since $(\sigma + j\omega\epsilon)$ and μ are ordinary scalar quantities and the order of a scalar product of vectors is interchangeable, the terms on the right-hand side cancel. If we integrate (5.69) over the volume of the waveguide junction enclosed by S shown in Fig. 5.1, we obtain

$$\int_S (\mathbf{E}^{(1)} \times \mathbf{H}^{(2)} - \mathbf{E}^{(2)} \times \mathbf{H}^{(1)}) \cdot \mathbf{n} dS = 0 \quad (5.70)$$

where use is made of Gauss's theorem to convert the volume integral to a surface integral. The contribution to the integral comes from that part of S consisting of the reference planes in the waveguides only. Let us now assume that $\mathbf{E}^{(1)}$ and $\mathbf{H}^{(1)}$ correspond to a set of currents which are all zero except $I_i^{(1)}$ at the i th port. Similarly, let $\mathbf{E}^{(2)}$ and $\mathbf{H}^{(2)}$ correspond to a set of currents which are zero except $I_j^{(2)}$ at the j th port. For the surface integral in (5.70), we then only have to consider the contributions from the i th and j th ports.

The integral of the first term in the bracket becomes

$$\begin{aligned}\int_S \mathbf{E}^{(1)} \times \mathbf{H}^{(2)} \cdot \mathbf{n} dS &= \int \mathbf{E}_{ij} V_j^{(1)} \times \mathbf{k} \times \mathbf{E}_{ij} I_j^{(2)} \cdot \mathbf{n} dS \\ &= V_j^{(1)} I_j^{(2)} \mathbf{k} \cdot \mathbf{n} = -V_j^{(1)} I_j^{(2)}\end{aligned}$$

where $\mathbf{E}_{ij}$ is the waveguide eigenfunction at the j th port, and the normalization condition for $\mathbf{E}_{ij}$ is used. The integral of the second term can be calculated in a similar manner, and (5.70) reduces to

$$-V_j^{(1)} I_j^{(2)} + V_i^{(2)} I_i^{(1)} = 0 \quad (5.71)$$

in this particular case. Since the I 's with superscript two are all zero except $I_j^{(2)}$, $V_i^{(2)}$ can be expressed by $Z_{ij} I_j^{(2)}$ from (5.64). Similarly, $V_j^{(1)}$ is given by $Z_{ji} I_i^{(1)}$. Substituting into (5.71), we have

$$(Z_{ij} - Z_{ji}) I_i^{(1)} I_j^{(2)} = 0$$

Noting that $I_i^{(1)}$ and $I_j^{(2)}$ can be specified arbitrarily, we must conclude

$$Z_{ij} = Z_{ji}$$

or equivalently,

$$\mathbf{Z} = \mathbf{Z}_t \quad (5.72)$$

The expression (5.72) is known as the reciprocal condition in ac circuit theory. The n -port network representing the waveguide junction is, therefore, reciprocal.

In much the same way, given the voltages instead of the currents at the ports, we obtain

$$\mathbf{Y} = \mathbf{Y}_t \quad (5.73)$$

Of course, if $\mathbf{Z}$ is nonsingular, (5.73) also follows from (5.68) and (5.72).

In the above discussion, we considered reference planes far from waveguide discontinuities; however, (5.64) has broader applications. For example, when there are several pairs of terminals inside the waveguide junction, the terminals in each pair being closely spaced compared to a wavelength $[\lambda = 2\pi/\omega(\epsilon\mu)^{1/2}]$, then the voltage between the terminals and the current flowing through them can be defined with little ambiguity. This follows because the electric field satisfies $\nabla \times \nabla \times \mathbf{E} - \omega^2 \epsilon\mu \mathbf{E} = 0$, assuming that σ is negligible, and that the first term on the left-hand side is made up of second derivatives which are individually much greater than the second term in the vicinity of the pair of terminals where field variation is large. As a result, the field should approximately satisfy $\nabla \times \nabla \times \mathbf{E} = 0$ which

is an equation for the static field. The voltage can be defined as the line integral of the electric field if the field resembles the static one. Some of the voltages and corresponding currents in (5.64) can, therefore, be considered to represent the conventional voltages and currents at the pairs of terminals. The linear relations and the reciprocal conditions (5.72) and (5.73) still hold due to the particular form of Maxwell's equations which have been used. However, a slight modification is required in the evaluation of the surface integral leading to (5.71). Let the j th port be the one representing such a pair of terminals, and let S_j be the closed surface enclosing a small volume in the vicinity of the pair of terminals as shown in Fig. 5.2. In order to evaluate the surface integral, we choose the elementary surface

$$\mathbf{n} dS = -\mathbf{e} d\xi \times \mathbf{h} d\eta$$

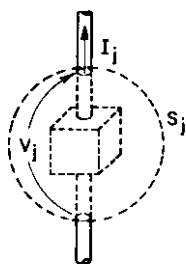


Fig. 5.2. A pair of terminals in junction and S_j .

where $\mathbf{e}$ and $\mathbf{h}$ are unit vectors in the direction of $\mathbf{E}^{(1)}$ and perpendicular to it, respectively. Then we have

$$\begin{aligned} \int_S \mathbf{E}^{(1)} \times \mathbf{H}^{(2)} \cdot \mathbf{n} dS &= - \int_{S_j} \mathbf{E}^{(1)} \times \mathbf{H}^{(2)} \cdot \mathbf{e} d\xi \times \mathbf{h} d\eta \\ &= - \iint_{S_j} \mathbf{E}^{(1)} \cdot \mathbf{e} d\xi \mathbf{H}^{(2)} \cdot \mathbf{h} d\eta = - V_j^{(1)} I_j^{(2)} \end{aligned}$$

from which (5.71) follows.

From the above discussion, we see that if a waveguide junction includes a small nonlinear device such as a semiconductor diode, then the junction can be treated as an n -port linear network with the device connected to one of the ports. This enables us to handle the nonlinearity of the device separately from the remainder of the microwave circuit. The n -port network in this case represents the junction which excludes a small volume occupied by the active nonlinear part of the device.

Now let us consider a linear transformation of the voltage and current given by

$$a_i = \frac{1}{2} |\operatorname{Re} Z_i|^{1/2} (V_i + Z_i I_i), \quad b_i = \frac{1}{2} |\operatorname{Re} Z_i|^{1/2} (V_i - Z_i^* I_i) \quad (5.74)$$

where $\operatorname{Re} Z_i$ is assumed to have a nonzero value. If the reference impedance Z_i is the impedance looking out from the i th port of the junction, a_i and b_i become the power waves introduced in Section 1.4. If Z_i is real and positive, and represents the characteristic impedance of the waveguide mode to which the i th port is assigned, then a_i and b_i are the traveling waves at the reference plane. Let $\mathbf{a}$ and $\mathbf{b}$ be the vectors with the i th component being a_i and b_i , respectively. If we define $\mathbf{F}$ and $\mathbf{G}$ by

$$\begin{aligned} \mathbf{F} &= \operatorname{diag} \left[\frac{1}{2} |\operatorname{Re} Z_1|^{-1/2} \quad \frac{1}{2} |\operatorname{Re} Z_2|^{-1/2} \quad \dots \quad \frac{1}{2} |\operatorname{Re} Z_n|^{-1/2} \right] \\ \mathbf{G} &= \operatorname{diag} [Z_1 \quad Z_2 \quad \dots \quad Z_n] \end{aligned} \quad (5.75)$$

we have

$$\mathbf{a} = \mathbf{F}(\mathbf{v} + \mathbf{G}\mathbf{i}), \quad \mathbf{b} = \mathbf{F}(\mathbf{v} - \mathbf{G}^+\mathbf{i}) \quad (5.76)$$

Since $\mathbf{v}$ and $\mathbf{i}$ are related linearly, there must be a linear relation between $\mathbf{a}$ and $\mathbf{b}$; let us write it in the form

$$\mathbf{b} = \mathbf{S}\mathbf{a} \quad (5.77)$$

where $\mathbf{S}$ is a square matrix. This $\mathbf{S}$ is called the scattering matrix. Substituting (5.76) into (5.77), we have

$$\mathbf{F}(\mathbf{v} - \mathbf{G}^+\mathbf{i}) = \mathbf{S}\mathbf{F}(\mathbf{v} + \mathbf{G}\mathbf{i})$$

Using $\mathbf{v} = \mathbf{Z}\mathbf{i}$, this reduces to

$$\mathbf{F}(\mathbf{Z} - \mathbf{G}^+)\mathbf{i} = \mathbf{S}\mathbf{F}(\mathbf{Z} + \mathbf{G})\mathbf{i}$$

from which $\mathbf{S}$ can be obtained in terms of $\mathbf{Z}$,

$$\mathbf{S} = \mathbf{F}(\mathbf{Z} - \mathbf{G}^+)(\mathbf{Z} + \mathbf{G})^{-1} \mathbf{F}^{-1} \quad (5.78)$$

Similarly, $\mathbf{Z}$ can be expressed in terms of $\mathbf{S}$,

$$\mathbf{Z} = \mathbf{F}^{-1}(\mathbf{I} - \mathbf{S})^{-1}(\mathbf{S}\mathbf{G} + \mathbf{G}^+)\mathbf{F} \quad (5.79)$$

The reciprocal condition of the junction in terms of $\mathbf{S}$ is given by

$$\mathbf{S}_t = \mathbf{P}\mathbf{S}\mathbf{P} \quad (5.80)$$

where

$$\mathbf{P} = \operatorname{diag} \left[\frac{\operatorname{Re} Z_1}{|\operatorname{Re} Z_1|} \quad \frac{\operatorname{Re} Z_2}{|\operatorname{Re} Z_2|} \quad \dots \quad \frac{\operatorname{Re} Z_n}{|\operatorname{Re} Z_n|} \right] \quad (5.81)$$

We shall now derive this formula from $\mathbf{Z} = \mathbf{Z}_t$. First, we note that $\mathbf{P}$ is a diagonal matrix whose i th diagonal component

$$p_i = \frac{\operatorname{Re} Z_i}{|\operatorname{Re} Z_i|} \quad (5.82)$$

is $+1$ or -1 depending on whether the real part of Z_i is positive or negative, respectively. Therefore, $\mathbf{P}$ is nonsingular, and $\mathbf{P} = \mathbf{P}^{-1}$. Substituting (5.79) into $\mathbf{Z} = \mathbf{Z}_t$, we have

$$\mathbf{F}^{-1}(\mathbf{I} - \mathbf{S})^{-1}(\mathbf{S}\mathbf{G} + \mathbf{G}^+) \mathbf{F} = \mathbf{F}(\mathbf{G}\mathbf{S}_t + \mathbf{G}^+)(\mathbf{I} - \mathbf{S}_t)^{-1} \mathbf{F}^{-1}$$

This is equivalent to

$$(\mathbf{S}\mathbf{G} + \mathbf{G}^+) \mathbf{F}^2 (\mathbf{I} - \mathbf{S}_t) = (\mathbf{I} - \mathbf{S}) \mathbf{F}^2 (\mathbf{G}\mathbf{S}_t + \mathbf{G}^+)$$

Expanding this and subtracting $\mathbf{G}^+ \mathbf{F}^2 - \mathbf{S}\mathbf{G}\mathbf{F}^2 \mathbf{S}_t$ from both sides, we have

$$\mathbf{S}\mathbf{G}\mathbf{F}^2 - \mathbf{G}^+ \mathbf{F}^2 \mathbf{S}_t = \mathbf{F}^2 \mathbf{G}\mathbf{S}_t - \mathbf{S}\mathbf{F}^2 \mathbf{G}^+$$

or equivalently,

$$\mathbf{S}\mathbf{F}^2 (\mathbf{G} + \mathbf{G}^+) = \mathbf{F}^2 (\mathbf{G} + \mathbf{G}^+) \mathbf{S}_t$$

However, since

$$\mathbf{F}^2 (\mathbf{G} + \mathbf{G}^+) = \frac{1}{2} \mathbf{P} = \frac{1}{2} \mathbf{P}^{-1}$$

the above equation gives (5.80), which is what we wished to derive.

When the signs of $\operatorname{Re} Z_i$ and $\operatorname{Re} Z_j$ are the same, (5.80) gives

$$S_{ij} = S_{ji}$$

and when they are opposite, it gives

$$S_{ij} = -S_{ji}$$

In either case, we have

$$|S_{ij}|^2 = |S_{ji}|^2 \quad (5.83)$$

Suppose that the a_i 's and the b_i 's represent power waves and none of the circuits contains a source, except the one connected to the i th port of the network. The power from the j th circuit to the network is given by $p_j(|a_j|^2 - |b_j|^2)$ and a_j ($j \neq i$) is equal to zero in this case. Therefore, the power to the j th circuit from the network is given by $p_j|b_j|^2$. Furthermore, b_j is equal to $S_{ji}a_i$ in the present case, and hence the ratio of the net power $p_j|b_j|^2$ into the j th circuit to the exchangeable power $p_i|a_i|^2$ from the i th circuit is equal to $p_i p_j |S_{ji}|^2$. Since (5.83) indicates that the value of this ratio does not change when the subscripts i and j are interchanged, we can

conclude that the relation between the net power into a load and the exchangeable power from the source remains the same in a reciprocal network when the roles of source and load are interchanged. This is a power reciprocal relation. Note in the above interchange of generator and load relation, only the location of the voltage source is changed while the impedances connected to the ports are kept constant.

When the a_i 's and b_i 's express traveling waves along the waveguides, (5.80) reduces to

$$\mathbf{S} = \mathbf{S}_t$$

since $\mathbf{P} = \mathbf{I}$ in this case. Although the relation (5.83) remains the same, the interpretation is different. Assume that the i th port is the only one with a nonzero incoming wave, then $|S_{ji}|^2$ is the ratio of the power emerging from the j th port to the incoming power at the i th port. This ratio remains constant when the roles of the i th and j th ports are interchanged. However, if the circuit connected to some port, say the k th port, is not matched to the characteristic impedance of the line, a reflection takes place resulting in an incoming wave, a_k which invalidates the assumption. Therefore, the power reciprocal relation $|S_{ij}|^2 = |S_{ji}|^2$ for traveling waves has only limited applications.

5.3 Lossless Conditions

When losses in a waveguide junction are small, it is convenient to consider a model in which the effect of losses is completely neglected. Let us consider the lossless conditions which $\mathbf{Z}$, $\mathbf{Y}$, and $\mathbf{S}$ must satisfy when these represent such an idealized model. Since the net power entering from the i th port is given by $\text{Re } V_i I_i^*$, the total power entering from all ports becomes

$$\sum \text{Re}(V_i I_i^*) = \sum \frac{1}{2}(V_i I_i^* + V_i^* I_i) = \frac{1}{2}(\mathbf{i}^+ \mathbf{v} + \mathbf{v}^+ \mathbf{i}) = \frac{1}{2} \mathbf{i}^+ (\mathbf{Z} + \mathbf{Z}^+) \mathbf{i} \quad (5.84)$$

where the summations are over all possible i 's and use is made of (5.66). For the lossless network, the total power must be zero regardless of the value of $\mathbf{i}$. Thus, we obtain

$$\mathbf{Z}^+ + \mathbf{Z} = 0 \quad (5.85)$$

It is obvious from the derivation that this is a necessary and sufficient condition for $\mathbf{Z}$ to express a lossless network. For one-port networks, (5.85) becomes $\text{Re } Z = 0$, which shows that the real part of the input impedance

is equal to zero. If we use (5.67) instead of (5.66), we obtain from (5.84)

$$\mathbf{Y} + \mathbf{Y}^+ = 0 \quad (5.86)$$

This is also necessary and sufficient for $\mathbf{Y}$ to represent a lossless network.

From (5.84), we find that $Z_{ij} = -Z_{ji}^*$ for a lossless network. Since $Z_{ji} = Z_{ij}$ for a reciprocal network, it follows that $Z_{ij} = -Z_{ij}^*$ for a lossless reciprocal network. This means that all the Z_{ij} 's are pure imaginary in this case. Similarly, all the Y_{ij} 's are pure imaginary for a lossless reciprocal network.

For a general passive network, since the total power expressed by (5.84) must not be negative, $\mathbf{Z} + \mathbf{Z}^+$ and $\mathbf{Y} + \mathbf{Y}^+$ have to be either positive-definite or semipositive-definite.

Next, let us consider the lossless condition in terms of $\mathbf{S}$. The net power into the i th port of a network is given by $p_i(|a_i|^2 - |b_i|^2)$. Therefore, the total power is given by

$$\sum p_i(|a_i|^2 - |b_i|^2) = \mathbf{a}^+ \mathbf{P} \mathbf{a} - \mathbf{b}^+ \mathbf{P} \mathbf{b} = \mathbf{a}^+ \mathbf{P} \mathbf{a} - \mathbf{a}^+ \mathbf{S}^+ \mathbf{P} \mathbf{S} \mathbf{a} = \mathbf{a}^+ (\mathbf{P} - \mathbf{S}^+ \mathbf{P} \mathbf{S}) \mathbf{a}$$

where use is made of (5.77). This must be zero regardless of the value of $\mathbf{a}$, and hence the lossless condition is given by

$$\mathbf{S}^+ \mathbf{P} \mathbf{S} = \mathbf{P} \quad (5.87)$$

When the $\text{Re} Z_i$'s are all positive, $\mathbf{P} = \mathbf{I}$, and (5.87) shows that $\mathbf{S}$ must be a unitary matrix.

Multiplying (5.87) by $\mathbf{P}$ from the right and utilizing (5.80), we find that $\mathbf{S}^+ \mathbf{S} = \mathbf{I}$ for a lossless reciprocal network. Taking the transposed matrix, this becomes

$$\mathbf{S} \mathbf{S}^* = \mathbf{I}$$

For a general passive network, since the total input power must not be negative, $\mathbf{P} - \mathbf{S}^+ \mathbf{P} \mathbf{S}$ has to be either positive-definite or semipositive-definite.

For a simple example, let us consider a lossless two-port network. Equation (5.87) gives three independent conditions

$$\begin{aligned} p_1 |S_{11}|^2 + p_2 |S_{21}|^2 &= p_1 \\ p_1 S_{11} S_{12}^* + p_2 S_{21} S_{22}^* &= 0 \\ p_1 |S_{12}|^2 + p_2 |S_{22}|^2 &= p_2 \end{aligned} \quad (5.88)$$

From the second condition, we have

$$|S_{11}|^2 |S_{12}|^2 = |S_{21}|^2 |S_{22}|^2$$

Combining the first and last conditions in (5.88) with this equation, we obtain

$$(p_2/p_1)(1 - |S_{22}|^2) |S_{11}|^2 = (p_1/p_2)(1 - |S_{11}|^2) |S_{22}|^2$$

which is equivalent to

$$|S_{11}|^2 = |S_{22}|^2$$

When S represents the power wave scattering matrix, this relation shows that the power reflection coefficient at one port is equal to that at the other port for a lossless two-port network. From this we conclude that the power reflection coefficient as well as the power transmission coefficient remains the same regardless of the position of the reference plane selected along a lossless transmission system. The fact that the exchangeable power is preserved during a nonsingular lossless transformation can also be easily shown using (5.88). Assume that a_2 is zero for the moment, then the exchangeable power from port 2 is given by

$$\frac{p_2 |b_2|^2}{1 - |S_{22}|^2} = \frac{p_2 |S_{21}|^2 |a_1|^2}{1 - |S_{11}|^2} = p_1 |a_1|^2$$

where we have used the first condition in (5.88). The right-hand side is exactly equal to the exchangeable power from the source connected to port 1. We have, therefore, shown that the exchangeable powers are the same at the input and output ports of a lossless two-port network, provided that $|S_{11}|^2 = |S_{22}|^2 \neq 1$. If $|S_{11}|^2 = |S_{22}|^2 = 1$, the input and output ports are effectively disconnected inside the network, and the transformation from the input port to the output port is said to be singular.

5.4 Frequency Characteristic of Lossless Reciprocal Junctions

Let us first consider Z and Y representing a lossless reciprocal junction. Maxwell's equations are given by

$$\nabla \times \mathbf{E} = -j\omega\mu\mathbf{H} \quad (5.89)$$

$$\nabla \times \mathbf{H} = j\omega\epsilon\mathbf{E} \quad (5.90)$$

Differentiating with respect to ω , we have

$$\nabla \times (\partial\mathbf{E}/\partial\omega) = -j\mu\mathbf{H} - j\omega\mu(\partial\mathbf{H}/\partial\omega) \quad (5.91)$$

$$\nabla \times (\partial\mathbf{H}/\partial\omega) = j\epsilon\mathbf{E} + j\omega\epsilon(\partial\mathbf{E}/\partial\omega) \quad (5.92)$$

Subtracting $(\partial \mathbf{E} / \partial \omega) \cdot$ times the complex conjugate of (5.90) from $\mathbf{H}^* \cdot$ (5.91), we have

$$\nabla \cdot \{(\partial \mathbf{E} / \partial \omega) \times \mathbf{H}^*\} = -j\mu \mathbf{H}^* \cdot \mathbf{H} - j\omega\mu (\partial \mathbf{H} / \partial \omega) \cdot \mathbf{H}^* + j\omega\epsilon \mathbf{E}^* \cdot (\partial \mathbf{E} / \partial \omega) \quad (5.93)$$

Similarly, we have from (5.89) and (5.92)

$$\nabla \cdot \{(\partial \mathbf{H} / \partial \omega) \times \mathbf{E}^*\} = j\epsilon \mathbf{E}^* \cdot \mathbf{E} + j\omega\epsilon (\partial \mathbf{E} / \partial \omega) \cdot \mathbf{E}^* - j\omega\mu \mathbf{H}^* \cdot (\partial \mathbf{H} / \partial \omega) \quad (5.94)$$

Subtracting (5.94) from (5.93), we obtain

$$\nabla \cdot \{(\partial \mathbf{E} / \partial \omega) \times \mathbf{H}^* - (\partial \mathbf{H} / \partial \omega) \times \mathbf{E}^*\} = -j(\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) \quad (5.95)$$

Let us integrate (5.95) throughout the volume within S in Fig. 5.1. The left integral becomes a surface integral over the waveguide cross section only, and is given by

$$\begin{aligned} & \int \{(\partial \mathbf{E} / \partial \omega) \times \mathbf{H}^* - (\partial \mathbf{H} / \partial \omega) \times \mathbf{E}^*\} \cdot \mathbf{n} \, dS \\ &= \sum \int \{ \mathbf{E}_{ii} (\partial V_{ii} / \partial \omega) \times (\mathbf{k} \times \mathbf{E}_{ii}) I_i^* - (\mathbf{k} \times \mathbf{E}_{ii}) (\partial I_i / \partial \omega) \times \mathbf{E}_{ii} V_i^* \} \cdot \mathbf{n} \, dS \\ &= - \sum \{ (\partial V_{ii} / \partial \omega) I_i^* + (\partial I_{ii} / \partial \omega) V_i^* \} = - \{ \mathbf{i}^+ (\partial v / \partial \omega) + \mathbf{v}^+ (\partial i / \partial \omega) \} \end{aligned}$$

where $\mathbf{E}_{ii}$ indicates the normalized eigenfunction of a waveguide propagating mode. The volume integral of (5.95), therefore gives

$$\{ \mathbf{i}^+ (\partial v / \partial \omega) + \mathbf{v}^+ (\partial i / \partial \omega) \} = j \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) \, dv \quad (5.96)$$

Using $\mathbf{v} = \mathbf{Z}\mathbf{i}$, the left-hand side is rewritten in the form

$$\begin{aligned} & \{ \mathbf{i}^+ (\partial \mathbf{Z} / \partial \omega) \mathbf{i} + \mathbf{i}^+ \mathbf{Z} (\partial \mathbf{i} / \partial \omega) + \mathbf{i}^+ \mathbf{Z}^+ (\partial \mathbf{i} / \partial \omega) \} \\ &= \{ \mathbf{i}^+ (\partial \mathbf{Z} / \partial \omega) \mathbf{i} + \mathbf{i}^+ (\mathbf{Z} + \mathbf{Z}^+) (\partial \mathbf{i} / \partial \omega) \} \end{aligned}$$

However, the second term on the right-hand side vanishes because of (5.85), and (5.96) therefore reduces to

$$\mathbf{i}^+ (\partial \mathbf{Z} / \partial \omega) \mathbf{i} = j \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) \, dv \quad (5.97)$$

For a lossless reciprocal junction, all the components of $\mathbf{Z}$ are pure imaginary, and if we write $\mathbf{Z} = j\mathbf{X}$, the components of $\mathbf{X}$ will be real. Substituting this

into (5.97), we have

$$\mathbf{i}^+ (\partial \mathbf{X} / \partial \omega) \mathbf{i} = \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv \quad (5.98)$$

This shows that $\partial \mathbf{X} / \partial \omega$ is positive definite, since the right-hand side is always positive for nonzero fields. If we use $\mathbf{i} = \mathbf{Y} \mathbf{v}$ instead of $\mathbf{v} = \mathbf{Z} \mathbf{i}$, we obtain from (5.96)

$$\mathbf{v}^+ (\partial \mathbf{B} / \partial \omega) \mathbf{v} = \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv \quad (5.99)$$

where $\mathbf{Y} = j\mathbf{B}$ is used. This shows that $\partial \mathbf{B} / \partial \omega$ is also positive definite.

For one-port networks, the above relations indicate that the derivatives with respect to ω of the input reactance X , and susceptance B , are always positive; i.e.,

$$(\partial X / \partial \omega) > 0, \quad (\partial B / \partial \omega) > 0$$

which correspond to Foster's reactance theorem. As illustrated in Fig. 5.3, X and B increase with ω almost everywhere.

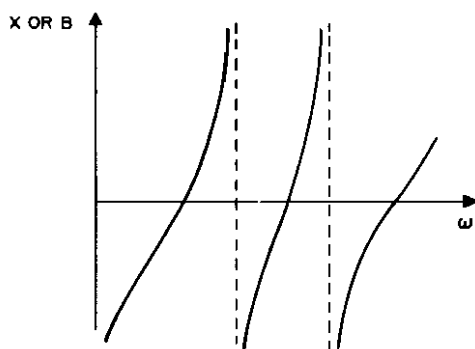


Fig. 5.3. Frequency characteristic of pure reactance X or pure susceptance B .

Next, let us consider the relation for $\mathbf{S}$ corresponding to (5.98) or (5.99). To do so, we express $\mathbf{v}$ and $\mathbf{i}$ in terms of $\mathbf{a}$ and $\mathbf{b}$ and substitute into (5.96). We have from (5.76)

$$\mathbf{a} - \mathbf{b} = \mathbf{F}(\mathbf{G} + \mathbf{G}^+) \mathbf{i}$$

Multiplying by $2\mathbf{F}\mathbf{P}$ from the left and using $2\mathbf{F}\mathbf{P}\mathbf{F}(\mathbf{G} + \mathbf{G}^+) = \mathbf{I}$, we obtain

$$\mathbf{i} = 2\mathbf{F}\mathbf{P}(\mathbf{a} - \mathbf{b}) \quad (5.100)$$

A combination of (5.76) and (5.100) gives

$$\mathbf{a} + \mathbf{b} = 2\mathbf{F}\mathbf{v} + \mathbf{F}(\mathbf{G} - \mathbf{G}^+) \mathbf{i} = 2\mathbf{F}\mathbf{v} + 2\mathbf{F}\mathbf{P}\mathbf{F}(\mathbf{G} - \mathbf{G}^+) (\mathbf{a} - \mathbf{b})$$

from which $\mathbf{v}$ can be obtained in terms of $\mathbf{a}$ and $\mathbf{b}$. If use is made of $(2\mathbf{F})^{-1} = \mathbf{P}\mathbf{F}(\mathbf{G} + \mathbf{G}^+)$, the expression for $\mathbf{v}$ is simplified as follows:

$$\begin{aligned} \mathbf{v} &= (2\mathbf{F})^{-1} (\mathbf{a} + \mathbf{b}) - \mathbf{P}\mathbf{F}(\mathbf{G} - \mathbf{G}^+) (\mathbf{a} - \mathbf{b}) \\ &= \mathbf{P}\mathbf{F} \{ (\mathbf{G} + \mathbf{G}^+) (\mathbf{a} + \mathbf{b}) - (\mathbf{G} - \mathbf{G}^+) (\mathbf{a} - \mathbf{b}) \} \\ &= 2\mathbf{P}\mathbf{F}(\mathbf{G}^+ \mathbf{a} + \mathbf{G}\mathbf{b}) \end{aligned} \quad (5.101)$$

Now, let us substitute (5.100) and (5.101) in the left-hand side of (5.96). After a little manipulation, we have

$$\begin{aligned} \{\mathbf{i}^+ (\partial \mathbf{v} / \partial \omega) + \mathbf{v}^+ (\partial \mathbf{i} / \partial \omega)\} &= 4 [\mathbf{a}^+ \{ \mathbf{F}^+ (\partial \mathbf{F} \mathbf{G}^+ / \partial \omega) + \mathbf{G} \mathbf{F}^+ (\partial \mathbf{F} / \partial \omega) \} \mathbf{a} \\ &\quad + \mathbf{a}^+ (\mathbf{F}^+ \mathbf{F} \mathbf{G}^+ + \mathbf{G} \mathbf{F}^+ \mathbf{F}) (\partial \mathbf{a} / \partial \omega) \\ &\quad - \mathbf{b}^+ \{ \mathbf{F}^+ (\partial \mathbf{F} \mathbf{G} / \partial \omega) + \mathbf{G}^+ \mathbf{F}^+ (\partial \mathbf{F} / \partial \omega) \} \mathbf{b} \\ &\quad - \mathbf{b}^+ (\mathbf{F}^+ \mathbf{F} \mathbf{G} + \mathbf{G}^+ \mathbf{F}^+ \mathbf{F}) (\partial \mathbf{b} / \partial \omega)] \end{aligned} \quad (5.102)$$

The second and fourth terms on the right-hand side can be simplified using the relations,

$$\mathbf{F}^+ \mathbf{F} \mathbf{G}^+ + \mathbf{G} \mathbf{F}^+ \mathbf{F} = \frac{1}{2} \mathbf{P}, \quad \mathbf{F}^+ \mathbf{F} \mathbf{G} + \mathbf{G}^+ \mathbf{F}^+ \mathbf{F} = \frac{1}{2} \mathbf{P} \quad (5.103)$$

The expressions within the brackets of the first and third terms on the right-hand side of (5.102) are more complicated. Let us calculate the i th diagonal component of the first one. Since the i th diagonal component of $\partial \mathbf{F} / \partial \omega$ is given by

$$\begin{aligned} \frac{1}{2} (\partial |\operatorname{Re} Z_i|^{-1/2} / \partial \omega) &= -\frac{1}{4} |\operatorname{Re} Z_i|^{-1/2} |\operatorname{Re} Z_i|^{-1} (\partial |\operatorname{Re} Z_i| / \partial \omega) \\ &= -\frac{1}{4} |\operatorname{Re} Z_i|^{-1/2} |Z_i^* + Z_i|^{-1} (\partial |Z_i^* + Z_i| / \partial \omega) \end{aligned}$$

the i th diagonal component of

$$\{\mathbf{F}^+ (\partial \mathbf{F} \mathbf{G}^+ / \partial \omega) + \mathbf{G} \mathbf{F}^+ (\partial \mathbf{F} / \partial \omega)\} = \mathbf{F}(\mathbf{G}^+ + \mathbf{G}) (\partial \mathbf{F} / \partial \omega) + \mathbf{F}^+ \mathbf{F} (\partial \mathbf{G}^+ / \partial \omega)$$

becomes

$$\begin{aligned} &-\frac{1}{8} |\operatorname{Re} Z_i|^{-1} \{ \partial (Z_i^* + Z_i) / \partial \omega \} + \frac{1}{4} |\operatorname{Re} Z_i|^{-1} (\partial Z_i^* / \partial \omega) \\ &= \frac{1}{4} |Z_i^* + Z_i|^{-1} \{ \partial (Z_i^* - Z_i) / \partial \omega \} \end{aligned}$$

Let $\mathbf{K}$ be a diagonal matrix whose i th component is given by

$$|Z_i^* + Z_i|^{-1} \{ \partial (Z_i^* - Z_i) / \partial \omega \}$$

then we have

$$\{\mathbf{F}^+ (\partial \mathbf{F} \mathbf{G}^+ / \partial \omega) + \mathbf{G} \mathbf{F}^+ (\partial \mathbf{F} / \partial \omega)\} = \frac{1}{4} \mathbf{K} \quad (5.104)$$

Noting that $\mathbf{K}$ is pure imaginary, the expression within the brackets of the third term on the right-hand side of (5.102) becomes

$$\{\mathbf{F}^+ (\partial \mathbf{F} \mathbf{G} / \partial \omega) + \mathbf{G}^+ \mathbf{F}^+ (\partial \mathbf{F} / \partial \omega)\} = \frac{1}{4} \mathbf{K}^+ = -\frac{1}{4} \mathbf{K} \quad (5.105)$$

Substituting (5.103), (5.104), and (5.105) into (5.102), we have

$$\{\mathbf{i}^+ (\partial \mathbf{v} / \partial \omega) + \mathbf{v}^+ (\partial \mathbf{i} / \partial \omega)\} = \mathbf{a}^+ \mathbf{K} \mathbf{a} + \mathbf{b}^+ \mathbf{K} \mathbf{b} + 2 \{\mathbf{a}^+ \mathbf{P} (\partial \mathbf{a} / \partial \omega) - \mathbf{b}^+ \mathbf{P} (\partial \mathbf{b} / \partial \omega)\} \quad (5.106)$$

Using $\mathbf{b} = \mathbf{S} \mathbf{a}$, the second term on the right-hand side becomes $\mathbf{a}^+ \mathbf{S}^+ \mathbf{K} \mathbf{S} \mathbf{a}$. Furthermore, the expression within the bracket of the last term can be rewritten in the form

$$\begin{aligned} \{\mathbf{a}^+ \mathbf{P} (\partial \mathbf{a} / \partial \omega) - \mathbf{b}^+ \mathbf{P} (\partial \mathbf{b} / \partial \omega)\} &= \mathbf{a}^+ \mathbf{P} (\partial \mathbf{a} / \partial \omega) - \mathbf{a}^+ \mathbf{S}^+ \mathbf{P} (\partial \mathbf{S} \mathbf{a} / \partial \omega) \\ &= \mathbf{a}^+ \mathbf{P} (\partial \mathbf{a} / \partial \omega) - \mathbf{a}^+ \mathbf{S}^+ \mathbf{P} (\partial \mathbf{S} / \partial \omega) \mathbf{a} \\ &\quad - \mathbf{a}^+ \mathbf{S}^+ \mathbf{P} \mathbf{S} (\partial \mathbf{a} / \partial \omega) \end{aligned} \quad (5.107)$$

The first and the last terms cancel each other because of the lossless condition (5.87), and only the second term remains. Consequently, (5.96) becomes

$$\mathbf{a}^+ \{\mathbf{K} + \mathbf{S}^+ \mathbf{K} \mathbf{S} - 2 \mathbf{S}^+ \mathbf{P} (\partial \mathbf{S} / \partial \omega)\} \mathbf{a} = j \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv$$

Dividing both sides by j , and using $-\mathbf{K} = \mathbf{K}^+$, we obtain

$$\mathbf{a}^+ j \{\mathbf{K}^+ + \mathbf{S}^+ \mathbf{K}^+ \mathbf{S} + 2 \mathbf{S}^+ \mathbf{P} (\partial \mathbf{S} / \partial \omega)\} \mathbf{a} = \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv \quad (5.108)$$

This shows that $j \{\mathbf{K}^+ + \mathbf{S}^+ \mathbf{K}^+ \mathbf{S} + 2 \mathbf{S}^+ \mathbf{P} (\partial \mathbf{S} / \partial \omega)\}$ is positive-definite.

When all the Z_i 's are real, $\mathbf{K} = 0$, and (5.108) reduces to

$$\mathbf{a}^+ 2j \mathbf{S}^+ \mathbf{P} (\partial \mathbf{S} / \partial \omega) \mathbf{a} = \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv \quad (5.109)$$

When the Z_i 's express the real characteristic impedances of the waveguide propagating modes, $\mathbf{P} = \mathbf{I}$, and we have

$$\mathbf{a}^+ 2j \mathbf{S}^+ (\partial \mathbf{S} / \partial \omega) \mathbf{a} = \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv$$

In order to see the physical meaning of (5.108), let us consider the simplest

case where a lossless reciprocal one-port network is connected to a generator having an internal impedance with a positive real part. Let the incident wave be the addition of two waves having equal amplitudes at frequencies ω and $\omega + \Delta\omega$

$$a(\omega) = Ae^{j\omega t} + Ae^{j(\omega + \Delta\omega)t}$$

Taking the real part and multiplying by $\sqrt{2}$, the actual wave as a function of time is given by

$$\begin{aligned} a(t) &= \sqrt{2}A \cos \omega t + \sqrt{2}A \cos(\omega + \Delta\omega)t \\ &= 2\sqrt{2}A \cos(\tfrac{1}{2}\Delta\omega t) \cos(\omega + \tfrac{1}{2}\Delta\omega)t \end{aligned}$$

The slowly varying function $2\sqrt{2}A \cos(\tfrac{1}{2}\Delta\omega t)$ expresses the envelope of the waveform.

We shall now consider the reflected wave. Since the magnitude of S_{11} is unity, from the lossless condition, S_{11} can be written in the form

$$S_{11} = e^{-j\varphi(\omega)} \quad (5.110)$$

If we write $\varphi(\omega + \Delta\omega) = \varphi + \Delta\varphi$, the reflected waves becomes

$$b(\omega) = Ae^{j(\omega t - \varphi)} + Ae^{j(\omega + \Delta\omega)t - (\varphi + \Delta\varphi)}$$

This represents an actual wave as a function of time which is given by

$$b(t) = 2\sqrt{2}A \cos\{\tfrac{1}{2}(\Delta\omega t - \Delta\varphi)\} \cos\{(\omega + \tfrac{1}{2}\Delta\omega)t - \varphi\}$$

The envelope of $b(t)$ lags behind that of $a(t)$ by $\Delta\varphi/\Delta\omega$. If no energy is to be transferred through the nodal points of the envelope, as we discussed in Section 3.2, the above result indicates that the energy is reflected back at a time $\Delta\varphi/\Delta\omega$ after its entrance into the network. Let τ be the limiting value of $\Delta\varphi/\Delta\omega$ when $\Delta\omega$ approaches zero, then τ becomes the time delay required for incident energy with angular frequency ω to enter the network and leave it again. Strictly speaking, $(2n\pi + \Delta\varphi)/\Delta\omega$ has to be considered as the time delay, where n is an integer since it is not certain which nodal points in the incident and reflected waves correspond to each other. However, if $n \neq 0$, $(2n\pi + \Delta\varphi)/\Delta\omega$ becomes infinite as $\Delta\omega$ approaches zero, and it cannot be considered as an appropriate value for the time delay. Therefore, n was chosen to be zero in the above discussion. From (5.109) and (5.110), τ is given by

$$\tau = (d\varphi/d\omega) = jS_{11}^{-1}(\partial S_{11}/\partial\omega) = \tfrac{1}{2}(a^*a)^{-1} \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv \quad (5.111)$$

The time delay is equal to the total stored energy in the network divided by the incident power. The larger the stored energy per unit incident power, the longer the delay.

5.5 Lossless Reciprocal Three-port Networks

To illustrate the utilization of some of the results obtained for the impedance matrix in Sections 5.2 and 5.3, let us derive an equivalent circuit for lossless, reciprocal, three-port networks. In general, a three-port network is characterized by a matrix of order 3×3 , and is therefore determined by nine complex numbers. However, since $Z_{ij} = Z_{ji}$ for reciprocal networks, only six complex numbers, or equivalently, twelve real numbers are sufficient to determine a reciprocal three-port network; furthermore, all the Z_{ij} 's are pure imaginary for lossless networks. Consequently, a lossless reciprocal three-port network is specified by only six real numbers. Let us choose these six real numbers to be $X_{11}, X_{12}, X_{13}, X_{22}, X_{23}$, and X_{33} where $X_{ij} = -jZ_{ij}$.

For the moment, we shall assume that X_{12}, X_{13} , and X_{23} do not vanish. Suppose that the transmission line connected to port III is open-circuited at a reference plane III' and the impedance looking out from port III becomes jX_3 . Then, V_3 is given by $-jX_3 I_3$. Substituting this into the simultaneous equations represented by $\mathbf{v} = \mathbf{Z}\mathbf{i}$, we can eliminate V_3 and I_3 . The result is given by

$$\begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = j \begin{bmatrix} \{X_{11} - X_{13}^2 (X_3 + X_{33})^{-1}\} & \{X_{12} - X_{13}X_{23} (X_3 + X_{33})^{-1}\} \\ \{X_{12} - X_{13}X_{23} (X_3 + X_{33})^{-1}\} & \{X_{22} - X_{23}^2 (X_3 + X_{33})^{-1}\} \end{bmatrix} \begin{bmatrix} I_1 \\ I_2 \end{bmatrix}$$

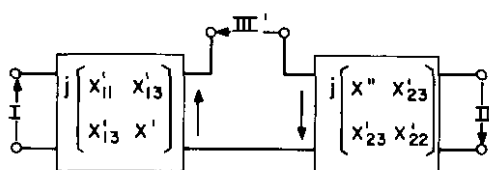
If we shift the open-circuited position III' along the transmission line, X_3 changes from $-\infty$ to $+\infty$; therefore, by choosing a proper position, we can make the off-diagonal components vanish, i.e.,

$$\{X_{12} - X_{13}X_{23}(X_3 + X_{33})^{-1}\} = 0 \quad (5.112)$$

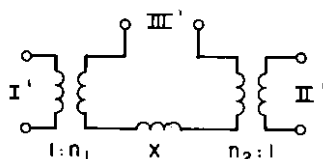
In other words, the coupling between ports I and II disappears when the transmission line is open-circuited at the proper position III'. If we look at the junction from ports I, II, and the new reference plane III', the impedance matrix must have the form

$$j \begin{bmatrix} X'_{11} & 0 & X'_{13} \\ 0 & X'_{22} & X'_{23} \\ X'_{13} & X'_{23} & X'_{33} \end{bmatrix}$$

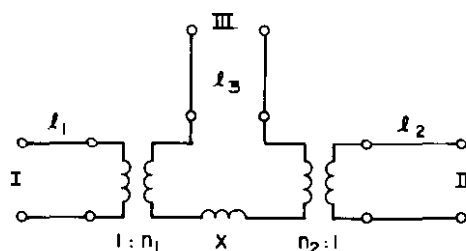
An equivalent circuit for this matrix is given by Fig. 5.4(a) where $X' + X'' =$



$$(a) \quad X' + X'' = X'_{33}$$



(b)



(c)

Fig. 5.4. Construction of an equivalent circuit for a lossless reciprocal three-port network.

X'_{33} . Since each of the two-port networks is lossless and reciprocal, they can each be represented by a series reactance, a transformer, and a certain length of transmission line, as shown from the discussion given in Section 1.3. The length of the line can always be chosen to be positive by adding an integer multiple of a wavelength. The equivalent circuit Fig. 5.4(a) can then be simplified to the form shown in Fig. 5.4(b) where the two series reactances are combined to give a single reactance X . From this discussion, an equivalent circuit for the junction looking in from the original ports I, II, and III can be represented by Fig. 5.4(c), where the line length attached to port III is again made positive by adding an integer multiple of a wavelength. In the equivalent circuit, the six real numbers characterizing the junction are l_1 , l_2 , l_3 , n_1 , n_2 , and X .

If either X_{23} or X_{13} is equal to zero, (5.112) may not be satisfied by any X_3 . In this case, we can use port I or port II in place of port III, in the above discussion, and the same type of equivalent circuit will be obtained. In case, any two of X_{12} , X_{13} , and X_{23} are equal to zero, say X_{12} and X_{13} , port I is completely isolated from ports II and III. The three-port network actually consists of a two-port network and an independent one-port network. Setting $n_1 = 0$, Fig. 5.4(c) can still serve as an equivalent circuit for the junction. Finally, if X_{12} , X_{13} , and X_{23} are all equal to zero, the junction consists of three independent one-port networks.

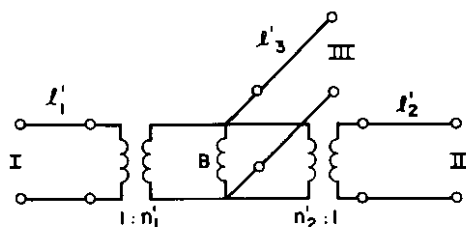


Fig. 5.5. Another equivalent circuit for a lossless reciprocal three-port network.

If we use the admittance matrix instead of the impedance matrix, an equivalent circuit, shown in Fig. 5.5, will be obtained, which is dual to Fig. 5.4(c).

5.6 Symmetrical Y Junctions

We shall now discuss some important properties of symmetrical Y junctions to familiarize ourselves with scattering matrices, as well as to prepare for Section 5.8 in which Y circulators will be studied. The circuit we shall discuss is a symmetrical junction with three identical transmission lines attached. These lines can be in the form of waveguides, coaxial lines, or any other type of transmission line supporting only one propagating mode. Let us assume that reference planes I, II, and III are symmetrically located with respect to the center of the junction, as shown in Fig. 5.6, and the reference impedances Z_i ($i = 1, 2, 3$) are all equal to each other. Since the reflection from a port corresponding to a unit incident wave must be the same from the symmetry regardless of the port number, we must have $S_{11} = S_{22} = S_{33}$. Similarly, the outgoing wave from port I, corresponding to a unit wave incident on port III, must be equal to the outgoing waves

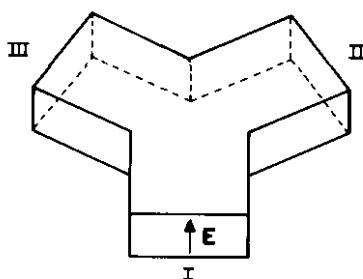


Fig. 5.6. Symmetrical Y junction.

from port III or port II corresponding to a unit wave incident on port II or port I, respectively, and we have $S_{13} = S_{32} = S_{21}$. In much the same way, we have another relation $S_{12} = S_{23} = S_{31}$. However, S_{13} may not be equal to S_{31} if no reciprocal condition is assumed. The most general form of the scattering matrix for the symmetrical Y junction is, therefore, given by

$$\mathbf{S} = \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{13} & S_{11} & S_{12} \\ S_{12} & S_{13} & S_{11} \end{bmatrix} \quad (5.113)$$

If identical incident waves are applied to all three ports simultaneously, the outgoing waves from all the ports are expected to be identical to each other from the symmetry. In other words, we expect that

$$\mathbf{x}_1 = \frac{1}{\sqrt{3}} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}$$

is an eigenvector for $\mathbf{S}$, where $\sqrt{3}$ is a normalizing factor. This can be verified by calculating $\mathbf{S}\mathbf{x}_1$ as follows.

$$\frac{1}{\sqrt{3}} \begin{bmatrix} S_{11} & S_{12} & S_{13} \\ S_{13} & S_{11} & S_{12} \\ S_{12} & S_{13} & S_{11} \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} = (S_{11} + S_{12} + S_{13}) \frac{1}{\sqrt{3}} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix}$$

Thus, $\mathbf{x}_1$ is, indeed, an eigenvector of $\mathbf{S}$, and the corresponding eigenvalue λ_1 is given by $S_{11} + S_{12} + S_{13}$. Similarly, if the incident waves have relative phase angles of 0° , 120° , and 240° , the outgoing waves will be expected to

have similar phase relations. Therefore,

$$\mathbf{x}_2 = \frac{1}{\sqrt{3}} \begin{bmatrix} 1 \\ e^{ja} \\ e^{j2a} \end{bmatrix}, \quad \mathbf{x}_3 = \frac{1}{\sqrt{3}} \begin{bmatrix} 1 \\ e^{j2a} \\ e^{ja} \end{bmatrix}$$

are expected to be eigenvectors of $\mathbf{S}$, where

$$a = (2\pi/3)$$

Since $\mathbf{S}\mathbf{x}_2$ and $\mathbf{S}\mathbf{x}_3$ are calculated to be

$$\mathbf{S}\mathbf{x}_2 = (S_{11} + S_{12}e^{ja} + S_{13}e^{j2a}) \mathbf{x}_2, \quad \mathbf{S}\mathbf{x}_3 = (S_{11} + S_{12}e^{j2a} + S_{13}e^{ja}) \mathbf{x}_3$$

$\mathbf{x}_2$ and $\mathbf{x}_3$ are, indeed, eigenvectors, and the corresponding eigenvalues are given by $\lambda_2 = S_{11} + S_{12}e^{ja} + S_{13}e^{j2a}$, and $\lambda_3 = S_{11} + S_{12}e^{j2a} + S_{13}e^{ja}$, respectively.

Let us consider the electric field along the axis of symmetry as the center of the junction. We choose a coordinate system with the z axis coinciding with the axis of symmetry, the y axis in the direction of the wave incident on port I, and the x axis perpendicular to the y and z axes as shown in Fig. 5.7. The field corresponding to $\mathbf{x}_1$ is given by the superposition of three cases:

- (i) Port I has an incident wave $1/\sqrt{3}$ and the remainder have none;
- (ii) Similarly, port II has $1/\sqrt{3}$ and the remainder none;
- (iii) Port III has $1/\sqrt{3}$ and the remainder none.

The axial component of the electric field in each case must be the same from the symmetry. Let us indicate it by E , then we have $3E$ at the center corresponding to $\mathbf{x}_1$. On the other hand, we have to superpose the following three cases for $\mathbf{x}_2$:

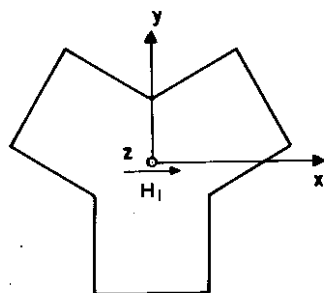


Fig. 5.7. Top view of symmetrical Y junction and the coordinate axes.

- (i) Port I has an incident wave $1/\sqrt{3}$ and the remainder have none;
- (ii) Port II has $e^{ja}/\sqrt{3}$ and the remainder none;
- (iii) Port III has $e^{j2a}/\sqrt{3}$ and the remainder none.

The axial components of electric field at the center are given by E for (i), Ee^{ja} for (ii) and Ee^{j2a} for (iii). Superposing these three fields, we have

$$E(1 + e^{ja} + e^{j2a}) = 0$$

which indicates that no axial component exists for $\mathbf{x}_2$. Similarly, the axial component corresponding to $\mathbf{x}_3$ is equal to zero.

Suppose that a thin conductor is introduced along the z -axis of the junction without destroying the symmetry; then the electromagnetic field corresponding to $\mathbf{x}_1$ is affected, but those corresponding to $\mathbf{x}_2$ and $\mathbf{x}_3$ are not. This conclusion follows from the fact that the thin conductor forces the electric field along it to vanish; whereas the axial components for $\mathbf{x}_2$ and $\mathbf{x}_3$ do not exist. Consequently, the eigenvalue λ_1 can be changed without affecting the eigenvalues λ_2 and λ_3 by the introduction of a thin conductor along the axis of symmetry.

Let us next consider the magnetic field perpendicular to the axis of symmetry. Let H_1 be the x component of the magnetic field when an incident wave of magnitude $1/\sqrt{3}$ is applied to port I; then the x components corresponding to waves $1/\sqrt{3}$ incident to port II and port III are given by $H_1 \cos a$ and $H_1 \cos 2a$, respectively. The corresponding y components are 0, $H_1 \sin a$ and $H_1 \sin 2a$. Suppose that the incident wave expressed by $\mathbf{x}_1$ is applied to the junction, then the x and y components of the magnetic field are obtained by the superposition of the three cases as follows:

$$H_x = H_1(1 + \cos a + \cos 2a) = 0, \quad H_y = H_1(0 + \sin a + \sin 2a) = 0$$

When $\mathbf{x}_2$ is applied, the same components are given by

$$\begin{aligned} H_x &= H_1(1 + e^{ja} \cos a + e^{j2a} \cos 2a) = \frac{3}{2}H_1 \\ H_y &= H_1(0 + e^{ja} \sin a + e^{j2a} \sin 2a) = \frac{3}{2}jH_1 \end{aligned}$$

To interpret this result, let us follow the usual procedure of multiplying the expressions by $\sqrt{2} e^{j\omega t}$ and taking the real parts which gives

$$H_x = \sqrt{2} |H_1| \frac{3}{2} \cos(\omega t + \theta), \quad H_y = -\sqrt{2} |H_1| \frac{3}{2} \sin(\omega t + \theta)$$

where $|H_1|$ and θ are defined through

$$H_1 = |H_1| e^{j\theta}$$

Drawing a diagram similar to Fig. 2.23, as shown in Fig. 5.8, we see that the magnetic field maintains a constant magnitude and rotates clockwise with angular velocity ω ; i.e., the field is circularly polarized. Similarly, corresponding to $\mathbf{x}_3$, we have

$$H_x = \frac{3}{2}H_1, \quad H_y = -\frac{3}{2}jH_1$$

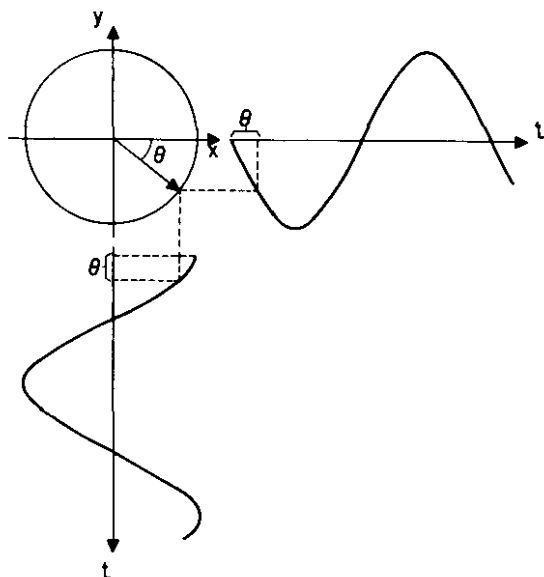


Fig. 5.8. Explanation for the circular polarization of magnetic field.

Therefore, the magnetic field for $\mathbf{x}_3$ is circularly polarized at the center and rotates counterclockwise with the same angular velocity ω as for $\mathbf{x}_1$. These results concerning the electric and magnetic fields at the center will be used in the discussion of Y circulators in Section 5.8.

Now, let us restrict ourselves to the symmetrical Y junctions satisfying the reciprocal condition discussed in Section 5.2. Since $\mathbf{P}$ is equal to either $\mathbf{I}$ or $-\mathbf{I}$ from the symmetry, (5.80) gives $S_{ij} = S_{ji}$. A comparison with (5.113) shows that $S_{13} = S_{12}$. The eigenvalues of S become

$$\lambda_1 = S_{11} + 2S_{12} \quad (5.114)$$

$$\lambda_2 = \lambda_3 = S_{11} - S_{12} \quad (5.115)$$

This indicates that $\mathbf{x}_2$ and $\mathbf{x}_3$ are degenerate.

Let us further assume that the junction is lossless. Then $\mathbf{S}$ must be a unitary matrix from (5.87); i.e.,

$$\mathbf{S}^+ \mathbf{S} = \mathbf{I}$$

Multiplying $\mathbf{x}_i (i=1, 2, 3)$ from the right and $\mathbf{x}_i^+$ from the left, we obtain

$$|\lambda_i|^2 = 1 \quad (i=1, 2, 3) \quad (5.116)$$

The magnitudes of the eigenvalues are, therefore, all unity. If we set $S_{11} = 0$, we have $\lambda_1 = 2S_{12}$ from (5.114), and $\lambda_2 = \lambda_3 = -S_{12}$ from (5.115). However, (5.116) cannot be satisfied simultaneously with these eigenvalues. We conclude from this that S_{11} cannot be equal to zero for any symmetrical *Y* junction if it is lossless and reciprocal. If all the diagonal components of a scattering matrix vanish, i.e., if all the ports are simultaneously matched, the junction is said to be totally matched. With this terminology, the above conclusion of $S_{11} \neq 0$ can be restated as follows: A symmetrical *Y* junction can never be totally matched if it is lossless and reciprocal. We can prove a slightly more general theorem: Any lossless reciprocal three-port network cannot be totally matched. Assume the contrary; then from the reciprocal condition (5.80) we have

$$(\mathbf{PS})_t = \mathbf{PS} = \begin{bmatrix} 0 & p_1 S_{12} & p_1 S_{13} \\ p_1 S_{12} & 0 & p_2 S_{23} \\ p_1 S_{13} & p_2 S_{23} & 0 \end{bmatrix}$$

which is symmetrical with respect to the main diagonal. From the lossless condition (5.87) or equivalently from $(\mathbf{PS})^+ \mathbf{P}(\mathbf{PS}) = \mathbf{P}$, we obtain

$$S_{13}^* S_{23} = 0, \quad S_{12}^* S_{23} = 0, \quad S_{12}^* S_{13} = 0$$

where the off-diagonal components of the above relation are calculated. To satisfy the three equations simultaneously, at least two of the components S_{12} , S_{13} , and S_{23} must be zero. Suppose that S_{12} and S_{13} are zero; then the first diagonal component of $(\mathbf{PS})^+ \mathbf{P}(\mathbf{PS}) = \mathbf{P}$ cannot be satisfied. Similarly, every possible combination of vanishing components leads to a contradiction showing that the initial assumption of total matching is wrong. This completes the proof.

Finally, let us ask how small we can make $|S_{11}|^2$ of a symmetrical *Y* junction when it is lossless and reciprocal? Eliminating S_{12} from (5.114) and (5.115), we have

$$3S_{11} = \lambda_1 + 2\lambda_2$$

Writing λ_1 and λ_2 in the forms

$$\lambda_1 = e^{j\theta_1}, \quad \lambda_2 = e^{j\theta_2}$$

the above equation gives

$$9|S_{11}|^2 = 5 + 4 \cos(\theta_1 - \theta_2)$$

Since the minimum value for $\cos(\theta_1 - \theta_2)$ is -1 , the minimum value of $|S_{11}|^2$ is given by

$$|S_{11}|^2 = \frac{1}{9}$$

In other words, one-ninth of the incident power is reflected back even with the best possible adjustment. The condition to give the minimum reflection is given by

$$\lambda_1 = -\lambda_2$$

As discussed previously, λ_1 can be changed by inserting a thin conductor along the axis of symmetry without affecting $\lambda_2 = \lambda_3$, and it is possible in most cases to satisfy the above condition.

5.7 Tensor Permeability of Ferrites

If ϵ , μ , and σ are ordinary scalar quantities, the product order in each term on the right-hand side of (5.69) can be changed without altering the value. Consequently, the terms on the right-hand side canceled each other, and the reciprocal relations were obtained. However, if the permeability is expressed by a matrix $[\mu]$ of order 3×3 and vectors are represented by matrices of order 3×1 , then $\mathbf{H}_2 \cdot [\mu] \mathbf{H}_1$ is not necessarily equal to $\mathbf{H}_1 \cdot [\mu] \mathbf{H}_2$ since the order of matrix product cannot generally be interchanged. The reciprocal relations do not follow, therefore, and the possibility arises that useful circuits can now be found which were not possible under the restriction of the reciprocal conditions, such as (5.80). For this reason, a large number of materials were investigated, and a class of magnetic materials called ferrites was experimentally found to have desired nonreciprocal characteristics when a static magnetic field was applied. The permeability of ferrites for small microwave fields is generally expressed in the form

$$[\mu] = \begin{bmatrix} \mu & -j\kappa & 0 \\ j\kappa & \mu & 0 \\ 0 & 0 & \mu_0 \end{bmatrix} \quad (5.117)$$

where the z direction is chosen to coincide with the internal static magnetic

field $\mathbf{H}_0$. The values of μ and κ in (5.117), and hence those of $\mu + \kappa$ and $\mu - \kappa$, change with $|\mathbf{H}_0|$. The typical variations are shown in Fig. 5.9.

The product of a matrix $[\mu]$ and a vector $\mathbf{H}$ gives another vector $\mathbf{B}$. A matrix which gives a vector in this way, when multiplied by another vector, is called a tensor, and ferrites are therefore said to have tensor permeabilities. The only information required for a discussion of microwave circuits

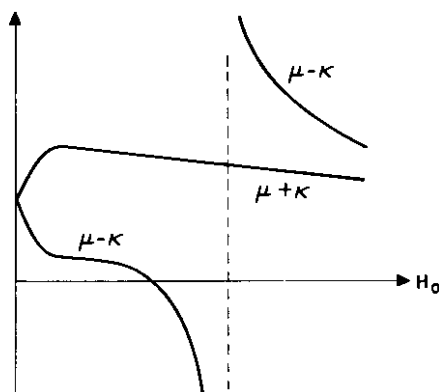


Fig. 5.9. The permeabilities $\mu + \kappa$ and $\mu - \kappa$ as function of static magnetic field.

containing ferrites is the knowledge concerning $[\mu]$ expressed in (5.117), and the variations of $\mu + \kappa$ and $\mu - \kappa$ with $|\mathbf{H}_0|$ as illustrated in Fig. 5.9. In this section we will briefly review the origin of such a tensor permeability.

Suppose that there is a small loop carrying a current I . Let A be the area inside the loop and $\mathbf{k}$ be the unit vector normal to the loop surface, as shown in Fig. 5.10. Then the current loop is said to have a magnetic moment given by $\mathbf{m} = I A \mathbf{k}$. In an atom each electron is orbiting and spinning, thereby

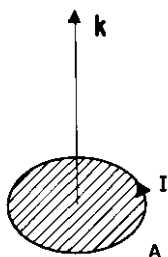


Fig. 5.10. Current loop.

forming two types of current loop. Ordinary substances can therefore be considered as a large number of current loops floating in a vacuum. Classifying the current loops by their current I and $A\mathbf{k}$, let us assume that there are many different kinds of current loops, and that the density of the j th kind is given by n_j per unit volume in the vicinity of $\mathbf{r}$. Summing all the the magnetic moments per unit volume, we can define a quantity $\mathbf{M}$ called the magnetization at $\mathbf{r}$ by

$$\sum_j I_j (A\mathbf{k})_j n_j = \mathbf{M} \quad (5.118)$$

The magnetization is the density of the total magnetic moment.

Now consider a macroscopic surface S in the space where the current loops exist, as shown in Fig. 5.11, and let us calculate the net current

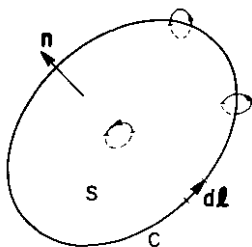


Fig. 5.11. Macroscopic surface S and its periphery C used in the derivation of $\mathbf{i}_m = \nabla \times \mathbf{M}$.

crossing the surface, which is given by $\int \mathbf{i}_m \cdot \mathbf{n} dS$, where the integration is over S ; $\mathbf{i}_m$ is the effective current density due to the current loops and $\mathbf{n}$ the unit vector normal to S . If C is the periphery of S , then the current loop which does not enclose C has no net contribution to the integral since, for such a loop, the same current crossing S from one side also crosses S from the other side. In order to calculate the net contribution, let us estimate the number of loops of the j th kind which encircle a short length dl of C . Although dl is considered to be macroscopically short and almost straight, it is very long compared to the size of the current loops. For a current loop to encircle dl , a certain point of the loop surface, say the center of gravity, has to be located inside a cylinder whose volume is $(A\mathbf{k})_j \cdot dl$, where dl is the vector representing the axis of the cylinder, as shown in Fig. 5.12. Since the density of the current loops is n_j , the number of current loops which encircle dl is given by $(A\mathbf{k})_j n_j \cdot dl$ on the average, each of which contributes I_j to the

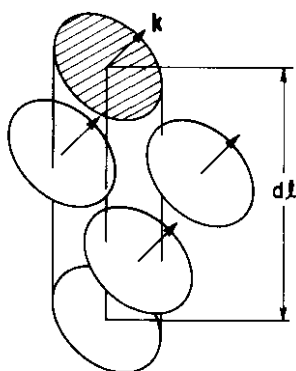


Fig. 5.12. Explanation of counting the number of current loops encircling dl .

integral. Summing all the contributions, we have

$$\int \mathbf{i}_m \cdot \mathbf{n} dS = \int_c \sum_j I_j (A\mathbf{k})_j n_j \cdot d\mathbf{l}$$

or equivalently

$$\int \mathbf{i}_m \cdot \mathbf{n} dS = \int_c \mathbf{M} \cdot d\mathbf{l}$$

The integral on the right-hand side can be rewritten using Stokes' theorem. After transposing to the left, the result gives

$$\int (\mathbf{i}_m - \nabla \times \mathbf{M}) \cdot \mathbf{n} dS = 0$$

Since this equation holds for any macroscopic surface S , we conclude that the effective current density is given by

$$\mathbf{i}_m = \nabla \times \mathbf{M} \quad (5.119)$$

whenever there is nonzero $\nabla \times \mathbf{M}$.

Maxwell's equation

$$\nabla \times \mathbf{H} = \mathbf{i} + (\partial \mathbf{D} / \partial t) \quad (5.120)$$

holds in magnetic materials as well as in vacuum where it can be written in the form

$$\nabla \times \mu_0^{-1} \mathbf{B} = \mathbf{i} + (\partial \mathbf{D} / \partial t)$$

However, this particular form is not valid in magnetic materials since $\mathbf{i}$

represents the free current density and does not include the effect of the current loops due to bounded electrons in atoms. If we include all the current effects on the right-hand side, it should hold in magnetic materials as well. Thus, we have

$$\nabla \times \mu_0^{-1} \mathbf{B} = \mathbf{i} + \mathbf{i}_m + (\partial \mathbf{D} / \partial t) = \mathbf{i} + \nabla \times \mathbf{M} + (\partial \mathbf{D} / \partial t)$$

or equivalently

$$\nabla \times (\mu_0^{-1} \mathbf{B} - \mathbf{M}) = \mathbf{i} + (\partial \mathbf{D} / \partial t) \quad (5.121)$$

A comparison of (5.120) and (5.121) shows that

$$\mathbf{B} = \mu_0 (\mathbf{H} + \mathbf{M}) \quad (5.122)$$

For a given applied $\mathbf{H}$, $\mathbf{B}$ becomes larger in magnetic materials than in vacuum due to the effects of bounded electrons in atoms which create $\mathbf{M}$ in (5.122).

In ordinary nonmagnetic materials, the effects of bounded electrons in atoms cancel each other and give no net contribution to $\mathbf{M}$. In ferrites, however, the magnetic moments due to the effective current loops of some of the spinning electrons remain uncanceled and give rise to the magnetic properties. An electron has its own mass and charge, and since the charge is negative, the effective current loop is formed in the opposite direction from the spin direction. Consequently, its magnetic moment $\mathbf{m}$ and angular momentum $\mathbf{J}$ have opposite directions as shown in Fig. 5.13. If we write the relation between $\mathbf{m}$ and $\mathbf{J}$ in the form

$$\mathbf{m} = \gamma \mathbf{J} \quad (5.123)$$

then γ must be negative. Also γ is considered to be one of the fundamental

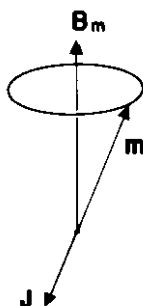


Fig. 5.13. Electron angular momentum $\mathbf{J}$ and magnetic moment $\mathbf{m}$.

constants of an electron, and its value has been experimentally determined as

$$\gamma = -1.76 \times 10^{11} \text{ [radian} \cdot \text{m}^2/\text{weber} \cdot \text{sec}]$$

When a current loop with magnetic moment $\mathbf{m}$ is placed in flux density $\mathbf{B}_m$, the loop receives a torque $\mathbf{m} \times \mathbf{B}_m$ in a direction which tends to maximize the flux linkage. This is the same kind of torque which causes a magnetic needle to point at the north pole. In general, the time derivative of angular momentum must be equal to the torque applied. In the case of an equivalent current loop due to an electron spinning about its axis in a ferrite, we have

$$\frac{d\mathbf{J}}{dt} = \mathbf{m} \times \mathbf{B}_m$$

Using (5.123), this becomes

$$\frac{d\mathbf{m}}{dt} = \gamma \mathbf{m} \times \mathbf{B}_m \quad (5.124)$$

As we can see from (5.122), $\mathbf{B}_m$ consists of two terms, one due to the internal magnetic field $\mathbf{H}_i$, and the other due to the magnetization $\mathbf{M}_i$. Let us now assume that the internal magnetic field is sufficiently strong to saturate the ferrite, in which case all the $\mathbf{m}$'s in a localized region will have the same direction. Then $\mu_0 \mathbf{M}_i$ gives no contribution to the right-hand side of (5.124), and $\mathbf{B}_m$ can be replaced by $\mu_0 \mathbf{H}_i$ since $\mathbf{m}$ and $\mathbf{M}_i$ are in the same direction. Multiplying by the number of $\mathbf{m}$'s per unit volume, (5.124) becomes

$$\frac{d\mathbf{M}_i}{dt} = \gamma \mu_0 \mathbf{M}_i \times \mathbf{H}_i \quad (5.125)$$

It is worth noting that $\mathbf{H}_i$ is generally not equal to the external magnetic field applied to the ferrite because of a demagnetizing effect.

Let us now write $\mathbf{H}_i$ as the sum of the static magnetic field $\mathbf{H}_0$ and the microwave field $\mathbf{H}e^{j\omega t}$, and let us write $\mathbf{M}_i$ as the sum of $\mathbf{M}_0$ and $\mathbf{M}e^{j\omega t}$:

$$\mathbf{H}_i = \mathbf{H}_0 + \text{Re} \{ \sqrt{2} \mathbf{H} e^{j\omega t} \}, \quad \mathbf{M}_i = \mathbf{M}_0 + \text{Re} \{ \sqrt{2} \mathbf{M} e^{j\omega t} \} \quad (5.126)$$

Assuming that the directions of $\mathbf{H}_0$ and $\mathbf{M}_0$ coincide, we take the z axis to be this direction. Furthermore, let us assume that the magnitudes of $\mathbf{H}$ and $\mathbf{M}$ are both small compared to those of $\mathbf{H}_0$ and $\mathbf{M}_0$, respectively. Substituting (5.126) into (5.125) and neglecting the product of $\mathbf{H}$ and $\mathbf{M}$, we have the following relations between the alternating components:

$$j\omega M_x = \gamma \mu_0 [M_y H_0 - M_0 H_y], \quad j\omega M_y = \gamma \mu_0 [M_0 H_x - M_x H_0], \quad j\omega M_z = 0$$

Solving these equations for M_x , M_y , and M_z , we have

$$\begin{aligned} M_x &= \frac{\gamma^2 \mu_0^2 M_0 H_0 H_x - j\omega\gamma\mu_0 M_0 H_y}{\gamma^2 \mu_0^2 H_0^2 - \omega^2} \\ M_y &= \frac{j\omega\gamma\mu_0 M_0 H_x + \gamma^2 \mu_0^2 M_0 H_0 H_y}{\gamma^2 \mu_0^2 H_0^2 - \omega^2} \\ M_z &= 0 \end{aligned} \quad (5.127)$$

The alternating component of magnetic induction $\mathbf{B}$ is given by $\mu_0 (\mathbf{H} + \mathbf{M})$; $\mathbf{B}$ can also be expressed as $[\mu] \mathbf{H}$, using the tensor permeability $[\mu]$ in (5.117); therefore, we have

$$\mathbf{B} = \mu_0 (\mathbf{H} + \mathbf{M}) = \begin{bmatrix} \mu & -j\kappa & 0 \\ j\kappa & \mu & 0 \\ 0 & 0 & \mu_0 \end{bmatrix} \mathbf{H}$$

A comparison of this equation with (5.127) shows that

$$\begin{aligned} \mu &= \mu_0 + \frac{\gamma^2 \mu_0^3 M_0 H_0}{\gamma^2 \mu_0^2 H_0^2 - \omega^2} = \mu_0 \left\{ 1 + \frac{\omega_0 \omega_M}{\omega_0^2 - \omega^2} \right\} \\ \kappa &= \frac{\omega\gamma\mu_0^2 M_0}{\gamma^2 \mu_0^2 H_0^2 - \omega^2} = -\mu_0 \frac{\omega\omega_M}{\omega_0^2 - \omega^2} \end{aligned} \quad (5.128)$$

where

$$\omega_0 = -\gamma\mu_0 H_0, \quad \omega_M = -\gamma\mu_0 M_0$$

The term ω_0 is the natural precession frequency of the electrons and is proportional to the magnitude of the static magnetic field H_0 . The term ω_M is proportional to M_0 , the saturation magnetization which is a property of the ferrite under consideration. From (5.128), we obtain

$$\mu + \kappa = \mu_0 \left\{ 1 + \frac{\omega_M}{\omega_0 + \omega} \right\}, \quad \mu - \kappa = \mu_0 \left\{ 1 + \frac{\omega_M}{\omega_0 - \omega} \right\} \quad (5.129)$$

Since ω_0 is proportional to H_0 , (5.129) gives a good explanation of why $\mu + \kappa$ and $\mu - \kappa$ vary with H_0 as shown in Fig. 5.9, excluding the region where H_0 is small, and the assumption of ferrite saturation is not valid.

When $\mathbf{H}$ is perpendicular to the z axis and circularly polarized, it is proportional to either

$$\begin{bmatrix} 1 \\ j \\ 0 \end{bmatrix} \quad \text{or} \quad \begin{bmatrix} 1 \\ -j \\ 0 \end{bmatrix}$$

depending on the direction of rotation as we discussed in Section 5.6. The magnetic induction $\mathbf{B}$ is given by $(\mu + \kappa) \mathbf{H}$ or $(\mu - \kappa) \mathbf{H}$ if $\mathbf{H}$ is proportional to the first or second vector, respectively. In other words, the effective permeability of the ferrite for the circularly polarized field becomes $\mu + \kappa$ or $\mu - \kappa$ depending on the direction of rotation.

Both μ and κ become infinite in (5.128) as ω approaches ω_0 ; however, in practice, various effects neglected in the above discussion, such as spin-spin and spin-lattice interactions, introduce losses, and consequently μ and κ remain finite. This situation is similar to that of a cavity near the resonant frequency; no matter how small the losses are, the fields are kept from growing infinite. In addition to the above-mentioned losses, introduced at or near ω_0 , when the applied magnetic field is below saturation, ferromagnetic domains are formed, and resonances due to the natural internal magnetic fields in the domains may also absorb microwave energy. Furthermore, the domains tend to move under the influence of low frequency microwave fields, thereby introducing another loss-mechanism in the ferrite. In order to minimize the effect of these "low-field" losses, H_0 must be made sufficiently large to saturate the ferrite.

5.8 Three-port Circulators

In Section 5.6, two of the eigenvalues of $\mathbf{S}$, λ_2 and λ_3 , were found to be equal to each other for a symmetrical Y junction when it was reciprocal. Suppose, however, that a piece of ferrite is introduced at the center of the Y junction without destroying the symmetry, then λ_2 and λ_3 may no longer be the same since the restricting reciprocal condition has been removed. The eigenvectors, on the other hand, remain the same because of the symmetry, and they are given by $\mathbf{x}_1$, $\mathbf{x}_2$, and $\mathbf{x}_3$ defined in Section 5.6.

Let us assume that λ_1 , λ_2 , and λ_3 are somehow adjusted to have the same magnitude but different phases 120° apart from each other; i.e.,

$$\lambda_2 = \lambda_1 e^{ja}, \quad \lambda_3 = \lambda_1 e^{j2a}$$

where a is given by $2\pi/3$, as before. Under this assumption, if a unit wave is incident on port I; i.e., if

$$\mathbf{a} = \begin{bmatrix} 1 \\ 0 \\ 0 \end{bmatrix} = (\mathbf{x}_1 + \mathbf{x}_2 + \mathbf{x}_3)/\sqrt{3}$$

is incident on the junction, then the reflected wave $\mathbf{b}$ is given by

$$\begin{aligned}\mathbf{b} &= \mathbf{S}\mathbf{a} = (\mathbf{S}\mathbf{x}_1 + \mathbf{S}\mathbf{x}_2 + \mathbf{S}\mathbf{x}_3)/\sqrt{3} \\ &= \lambda_1 (\mathbf{x}_1 + e^{j\theta}\mathbf{x}_2 + e^{j2\theta}\mathbf{x}_3)/\sqrt{3} \\ &= \lambda_1 \begin{bmatrix} 0 \\ 0 \\ 1 \end{bmatrix}\end{aligned}$$

where use is made of the fact that the eigenvectors are multiplied by their eigenvalues when $\mathbf{S}$ is applied. The above result indicates that λ_1 emerges from port III, and nothing emerges from port II when a unit wave is incident on port I. Similarly, when a unit wave is incident on port II or port III, λ_1 emerges from port I or port II, respectively. The junction is called a circulator when there is a rotational relation between the incident and reflected waves as illustrated in this example. If the order of leading phases is changed from λ_1, λ_2 , and λ_3 to λ_1, λ_3 , and λ_2 , an incident wave on ports I, II, and III emerges from ports II, III, and I, respectively, and a circulator with rotational direction opposite to the previous case will be obtained.

With this much preparation, let us now consider how the phases of the eigenvalues can be adjusted to 120° apart. Let $\mathbf{v}_i$ and $\mathbf{i}_i$ be the voltage and current at the junction ports when $\mathbf{x}_i$ is incident. Furthermore, let the superscript (0) indicate the case with no static magnetic field applied to the ferrite. Noting that $\mathbf{b} = \mathbf{S}\mathbf{x}_i = \lambda_i\mathbf{x}_i$, we have from (5.100) and (5.101)

$$\begin{aligned}(\mathbf{v}_i^{(0)})_t \mathbf{i}_j - (\mathbf{v}_j)_t \mathbf{i}_i^{(0)} &= (\mathbf{G}^+ \mathbf{x}_i + \mathbf{G}\lambda_i^{(0)}\mathbf{x}_i)_t 2\mathbf{F}_t \mathbf{P}_t 2\mathbf{F}\mathbf{P}(\mathbf{x}_j - \lambda_j\mathbf{x}_j) \\ &\quad - (\mathbf{G}^+ \mathbf{x}_j + \mathbf{G}\lambda_j\mathbf{x}_j)_t 2\mathbf{F}_t \mathbf{P}_t 2\mathbf{F}\mathbf{P}(\mathbf{x}_i - \lambda_i^{(0)}\mathbf{x}_i)\end{aligned}$$

Since $\mathbf{F}$, $\mathbf{G}$, and $\mathbf{P}$ are diagonal matrices, transposing them introduces no change, and the order of their product can be interchanged. Furthermore, since each term on the right-hand side is a matrix of order 1×1 , we can take the transposed matrix of the second term without changing the result. Thus, a little manipulation shows that

$$(\mathbf{v}_i^{(0)})_t \mathbf{i}_j - (\mathbf{v}_j)_t \mathbf{i}_i^{(0)} = (\mathbf{x}_i)_t (\lambda_i^{(0)} - \lambda_j) 4\mathbf{F}^2 (\mathbf{G} + \mathbf{G}^+) \mathbf{x}_j$$

Using the relation $4\mathbf{F}^2 (\mathbf{G} + \mathbf{G}^+) = 2\mathbf{P}$, this reduces to

$$(\mathbf{v}_i^{(0)})_t \mathbf{i}_j - (\mathbf{v}_j)_t \mathbf{i}_i^{(0)} = 2(\lambda_i^{(0)} - \lambda_j) (\mathbf{x}_i)_t \mathbf{P}\mathbf{x}_j \quad (5.130)$$

The left-hand side can be rewritten in terms of the electric and magnetic

fields at the port reference planes:

$$(\mathbf{v}_i^{(0)})_i \mathbf{i}_j - (\mathbf{v}_j)_i \mathbf{i}_i^{(0)} = \sum \int \{ \mathbf{E}_i^{(0)} \times \mathbf{H}_j - \mathbf{E}_j \times \mathbf{H}_i^{(0)} \} \cdot \mathbf{k} dS$$

where the summation is over all the ports. Let $\mathbf{n}$ be the outward unit vector normal to a closed surface S enclosing the junction; then $\mathbf{n} = -\mathbf{k}$ over the reference planes, as illustrated in Fig. 5.1. Therefore, the right-hand side becomes

$$- \sum \int \{ \mathbf{E}_i^{(0)} \times \mathbf{H}_j - \mathbf{E}_j \times \mathbf{H}_i^{(0)} \} \cdot \mathbf{n} dS = - \int \mathbf{V} \cdot \{ \mathbf{E}_i^{(0)} \times \mathbf{H}_j - \mathbf{E}_j \times \mathbf{H}_i^{(0)} \} dv$$

where use is made of Gauss's theorem. Following the method employed in (5.69), the volume integral can now be calculated

$$\int \mathbf{V} \cdot \{ \mathbf{E}_i^{(0)} \times \mathbf{H}_j - \mathbf{E}_j \times \mathbf{H}_i^{(0)} \} dv = \int j\omega \mathbf{H}_i^{(0)} \cdot ([\mu] - \mu^{(0)}) \mathbf{H}_j dv$$

As a result, (5.130) becomes

$$(\lambda_j - \lambda_i^{(0)}) (\mathbf{x}_i)_i \mathbf{P} \mathbf{x}_j = (j\omega/2) \int \mathbf{H}_i^{(0)} \cdot ([\mu] - \mu^{(0)}) \mathbf{H}_j dv \quad (5.131)$$

Setting $i = j = 1$, the change of λ_1 , due to the application of a static magnetic field on the ferrite, is calculated to be

$$\lambda_1 - \lambda_1^{(0)} = \pm \frac{1}{2} j\omega \int \mathbf{H}_1^{(0)} \cdot ([\mu] - \mu^{(0)}) \mathbf{H}_1 dv$$

where the upper and lower signs correspond to cases where $\mathbf{P}$ is equal to $\mathbf{I}$ and $-\mathbf{I}$, respectively. Similarly, setting either $i = 3, j = 2$ or $i = 2, j = 3$, we have

$$\lambda_2 - \lambda_2^{(0)} = \pm \frac{1}{2} j\omega \int \mathbf{H}_3^{(0)} \cdot ([\mu] - \mu^{(0)}) \mathbf{H}_2 dv$$

$$\lambda_3 - \lambda_3^{(0)} = \pm \frac{1}{2} j\omega \int \mathbf{H}_2^{(0)} \cdot ([\mu] - \mu^{(0)}) \mathbf{H}_3 dv$$

where use is made of the relation $\lambda_2^{(0)} = \lambda_3^{(0)}$.

Now suppose that the static magnetic field is gradually increased in the axial direction of the Y junction. The initial contribution to the integral in the first equation from the axial component of $\mathbf{H}_1$ is negligibly small, since $\mu_0 = \mu^{(0)}$. Furthermore, the component perpendicular to the axis is zero at the center and small in the vicinity of the center; hence, $|\lambda_1 - \lambda_1^{(0)}|$ will

remain small. On the other hand, $\mathbf{H}_2$ and $\mathbf{H}_3$ have perpendicular components, and both $|\lambda_2 - \lambda_2^{(0)}|$ and $|\lambda_3 - \lambda_3^{(0)}|$ grow rapidly with increasing static magnetic field. The magnetic fields $\mathbf{H}_2$ and $\mathbf{H}_3$ are respectively proportional to

$$\begin{bmatrix} 1 \\ j \\ 0 \end{bmatrix} \quad \text{and} \quad \begin{bmatrix} 1 \\ -j \\ 0 \end{bmatrix}$$

along the center axis; consequently, we have

$$\begin{aligned} \mathbf{H}_3^{(0)} \cdot [\mu] \mathbf{H}_2 &\simeq H_3^{(0)} H_2 \begin{bmatrix} 1 & -j & 0 \end{bmatrix} \begin{bmatrix} \mu & -j\kappa & 0 \\ j\kappa & \mu & 0 \\ 0 & 0 & \mu_0 \end{bmatrix} \begin{bmatrix} 1 \\ j \\ 0 \end{bmatrix} \\ &= H_3^{(0)} H_2 2(\mu + \kappa) \end{aligned}$$

where $H_3^{(0)}$ and H_2 are complex quantities representing the x components of $\mathbf{H}_3^{(0)}$ and $\mathbf{H}_2$, respectively. Similarly, we have

$$\mathbf{H}_2^{(0)} \cdot [\mu] \mathbf{H}_3 \simeq H_2^{(0)} H_3 2(\mu - \kappa)$$

From these relations, the changes of λ_2 and λ_3 due to the static magnetic field are given by

$$\lambda_2 - \lambda_2^{(0)} \simeq \pm j\omega \{(\mu + \kappa) - \mu^{(0)}\} \int H_3^{(0)} H_2 dv$$

$$\lambda_3 - \lambda_3^{(0)} \simeq \pm j\omega \{(\mu - \kappa) - \mu^{(0)}\} \int H_2^{(0)} H_3 dv$$

When the static magnetic field H_0 is small, μ and κ vary with H_0 , as shown in Fig. 5.14, and μ remains approximately equal to $\mu^{(0)}$. Furthermore, $H_2^{(0)}$ is equal to $H_3^{(0)}$ and, both H_2 and H_3 are expected to be initially close to

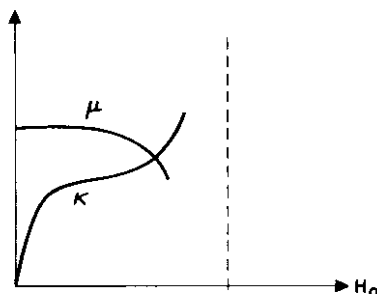


Fig. 5.14. The permeabilities μ and κ versus static magnetic field H_0 .

this value which corresponds to zero static magnetic field. Therefore, the above equations show that $\lambda_2 - \lambda_2^{(0)}$ and $\lambda_3 - \lambda_3^{(0)}$ must change proportionally to κ and $-\kappa$, respectively. That is, with increasing H_0 , λ_2 and λ_3 move in opposite directions starting from the original value $\lambda_2^{(0)} = \lambda_3^{(0)}$. If we neglect the losses in the junction, the magnitudes of the eigenvalues become unity, and they should appear as illustrated in Fig. 5.15(a). If we further increase

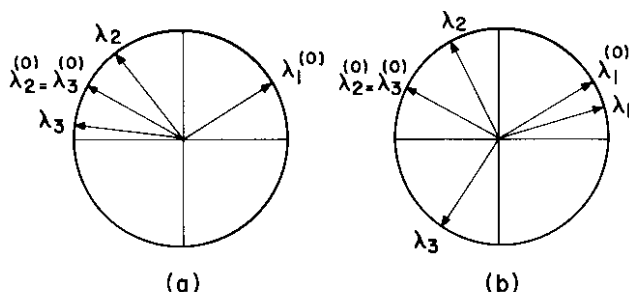


Fig. 5.15. Variations of the eigenvalues λ_1 , λ_2 , and λ_3 with static magnetic field.

H_0 , then H_2 and H_3 may no longer be close to their original value, and hence λ_2 and λ_3 cease to be symmetrical with respect to the original eigenvalue $\lambda_2^{(0)} = \lambda_3^{(0)}$. At some value of H_0 , the angle between λ_2 and λ_3 will become 120° , and λ_1 may have shifted slightly from $\lambda_1^{(0)}$ since $\mathbf{H}_1$ is not exactly equal to zero over the ferrite. As a result, λ_1 , λ_2 , and λ_3 will have taken the positions in Fig. 5.15(b). If we now insert a thin conductor along the center axis and adjust λ_1 , without changing λ_2 and λ_3 as explained in Section 5.6, until the angles from λ_1 to λ_2 and λ_3 become 120° , then the junction becomes a circulator.

In practice, we carry out the reverse process to minimize the insertion loss. First, we choose the dimension of the ferrite in such a way that it resonates at the desired frequency as a transmission type cavity (i.e., no power transmission takes place at detuned frequencies) with no static magnetic field applied. To observe the resonance clearly, we may have to decrease the coupling between the waveguides and the central part of the junction where the ferrite is located. Once the dimensions are properly chosen, we apply the static magnetic field and observe the split of the resonance curve due to the removal of the degeneracy. If a wrong resonant mode was picked up, the split would not occur. As we increase the static magnetic field, the separation between the two resonances becomes wider, and at a certain point, the lower

resonant frequency becomes stationary and then starts increasing. This corresponds to the point where $\mu + \kappa$ starts decreasing in Fig. 5.9 and indicates that the ferrite is saturated. We fix the static magnetic field at this value to obtain the minimum low field losses without suffering excessive losses due to the ferromagnetic resonance discussed at the end of Section 5.7. We then adjust the coupling between the waveguides and the central part of the junction to obtain proper circulator action. The reason for the success of this procedure can be seen as follows. At detuned frequencies, the eigenvalues λ_1 , λ_2 , and λ_3 are identical if the original resonance is a transmission type. As we increase the frequency and pass through the resonances, λ_2 first rotates 360° , since $\mathbf{H}_2$ sees $\mu + \kappa$, and then λ_3 rotates 360° , while λ_1 more or less stands still. By changing the couplings, we change the Q_{ext} 's, and hence the speed with which λ_2 and λ_3 rotate. When the speed is properly adjusted, λ_1 , λ_2 , and λ_3 will be positioned 120° apart from each other at (or at least near) the desired frequency, and the junction will become a circulator.

Let Z_0 be the reference impedance for the symmetrical circulator, then the impedance looking into the circulator from each port must be Z_0^* . If we connect a lossless two-port network to port I in order to transform Z_0^* to Z_1^* , and consider the other port of the two-port network as a new port of the junction with reference impedance Z_1 , then the junction acts as a circulator with the reference impedances Z_1 , Z_0 , and Z_0 at ports I, II, and III, respectively. Similarly, by adding a lossless two-port network which transforms Z_0^* to Z_2^* at port II and another which transforms Z_0^* to Z_3^* at port III, we obtain a circulator with reference impedances Z_1 , Z_2 , and, Z_3 ; these impedances are arbitrary as long as their real parts have the same sign as that of Z_0 .

The general form of a lossless circulation is given by either

$$\mathbf{S} = \begin{bmatrix} 0 & e^{j\theta_1} & 0 \\ 0 & 0 & e^{j\theta_2} \\ e^{j\theta_3} & 0 & 0 \end{bmatrix} \quad \text{or} \quad \mathbf{S} = \begin{bmatrix} 0 & 0 & e^{j\theta_1} \\ e^{j\theta_2} & 0 & 0 \\ 0 & e^{j\theta_3} & 0 \end{bmatrix}$$

depending on the direction of the circulation, where θ_1 , θ_2 , and θ_3 are arbitrary real angles. Now suppose that the scattering matrix is given by the first form, and that an impedance Z , instead of Z_3 , is connected to port III, and a unit power wave is incident on port I. Since the relation between the voltage and current at port III is given by $V_3 = -ZI_3$, a_3 is given by

$$a_3 = \{(Z - Z_3)/(Z + Z_3^*)\} b_3 \quad (5.132)$$

where use is made of (5.74) with $i=3$. Substituting this into the relation $\mathbf{b} = \mathbf{S}\mathbf{a}$, and noting that $a_1 = 1$ and $a_2 = 0$, we have

$$b_2 = e^{j(\theta_2 + \theta_3)} \{(Z - Z_3)/(Z + Z_3^*)\}$$

Assuming that the real parts of the Z_i 's are all positive, the output power at port II therefore becomes $|(Z - Z_3)/(Z + Z_3^*)|^2$ times the incident power on port I. If $\text{Re} Z$ is negative, $|(Z - Z_3)/(Z + Z_3^*)| > 1$ and power amplification is obtained.

PROBLEMS

- 5.1 Following the discussion on self-adjoint matrices, obtain the spectral representation of a unitary matrix. Show also that if two unitary matrices commute, they can be simultaneously diagonalized by a unitary matrix through the similarity transformation.
- 5.2 Suppose that three rectangular waveguides are connected to form a four-port junction, as shown in Fig. 5.16(a), and each has only one propagating mode TE_{10} ,

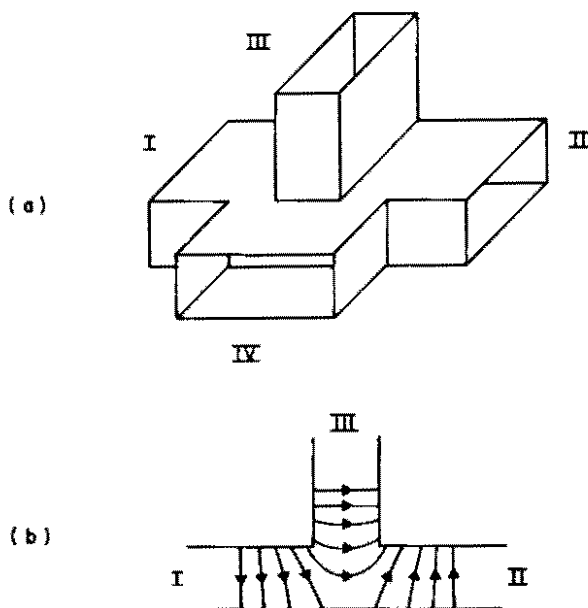


Fig. 5.16. (a) Magic T, (b) Intuitive explanation of the phase relation between ports I and II.

then an incident wave on port III will be divided in two and will emerge from ports I and II with equal, but out-of-phase, amplitudes, as shown in Fig. 5.16(b). Port IV has no output since the wave into port III tries to excite a TE_{01} mode in the port IV waveguide which is in the cutoff region. Similarly, an incident wave on port IV will emerge from ports I and II in-phase but nothing emerges from port III. This waveguide junction is called a magic T . Assuming total matching, calculate the scattering matrix of a lossless magic T .

- 5.3 Repeat the discussion given in Section 5.6 for symmetrical cross junctions.
- 5.4 Calculate the internal static magnetic field H_0 that is necessary for the natural precession frequency of electrons to be 4000 MHz.
- 5.5 Suppose that a thin ferrite rod is inserted along the center axis of a circular waveguide and a static magnetic field is applied in the axial direction. Then the plane of polarization of the TE_{11} mode will rotate with distance along the axis. This phenomenon is known as Faraday rotation. By decomposing the linearly polarized wave into two oppositely rotating circularly polarized waves, explain why such a rotation takes place.
- 5.6 Adjust the length of the above ferrite rod and/or the static magnetic field so as to rotate the plane of polarization 45° and connect transducers from the circular TE_{11} to the rectangular TE_{10} mode at both ends, after twisting one 45° with respect to the other, as shown in Fig. 5.17. If a resistive film is placed in each transducer to

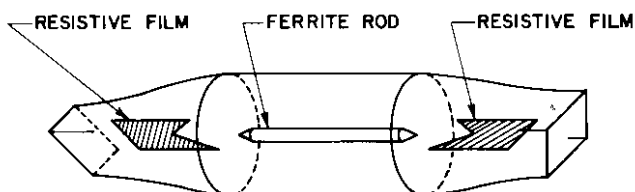


Fig. 5.17. Faraday rotation isolator.

absorb the mode corresponding to the rectangular TE_{01} , the device acts as an isolator; the wave passes through the device in one direction with little attenuation, whereas it is almost completely absorbed in the other direction. Explain qualitatively how the isolator action takes place.

- 5.7 Prove that a lossless symmetrical Y junction is a circulator if it is totally matched.
- 5.8 If the ferrite dimension for a symmetrical three port circulator is chosen so as to reflect the incident power at the resonant frequency, and if a separate resonator is properly coupled to x_1 mode only, then, with a proper magnetic field applied, the circulator action takes place over a relatively wide frequency range. This is because all three vectors λ_1 , λ_2 , and λ_3 rotate in the same direction with similar velocities maintaining the proper phase relation (120° apart) near the resonance. Explain qualitatively how this proper phase relation can be obtained.

- 5.9 Prove that, if reference impedances are changed from Z_i to Z_i' ($i = 1, 2, \dots, n$), the scattering matrix is transformed from $\mathbf{S}$ to $\mathbf{S}'$;

$$\mathbf{S}' = \mathbf{A}^{-1}(\mathbf{S} - \mathbf{\Gamma}^*) (\mathbf{I} - \mathbf{\Gamma}\mathbf{S})^{-1}\mathbf{A}^*$$

where $\mathbf{\Gamma}$ and $\mathbf{A}$ are the diagonal matrices with i th diagonal components $r_i = (Z_i' - Z_i) / (Z_i' + Z_i^*)$ and $|1 - r_i r_i^*|^{1/2} (1 - r_i^*) / |1 - r_i|$, respectively.



CHAPTER 6

COUPLED MODES AND PERIODIC STRUCTURES

When several transmission lines are coupled together through perturbations of various kinds, we can treat the resultant complicated transmission system as a waveguide and develop an appropriate new theory for it as we did in Chapter 3. Instead of proceeding in this manner, however, it is more advantageous in many cases to discuss the complicated system as a combination of individual simple transmission lines; in this way we can make full use of knowledge concerning individual transmission line properties since these are generally well understood from separate studies. In this chapter, we shall study one such approach called the theory of coupled modes; this provides a powerful tool for the understanding of relatively complicated combined systems often encountered in practice.

The discussion will reduce to an eigenvalue problem of a matrix. However, since the matrix is not self-adjoint, a new method is developed to study the eigenvalue problem. This method is applicable to many problems, including the discussion of periodic structures.

First, we set up an equation for a system with lossless distributed coupling, and then deduce the eigenvalue problem. After discussing the general properties of the solutions, we solve the eigenvalue problem using a perturbation method. We then discuss in detail the interaction of two waves as the simplest, important example of coupled modes. We next formulate two eigenvalue problems for periodic structures, one for lossless systems, and the other for reciprocal systems, and then discuss the properties of their solutions following the method developed for distributed coupling. Finally, taking the specific example of a waveguide with periodically placed identical

discontinuities, we discuss the ω - β diagram and the space harmonics; these are found to be particularly useful concepts when a periodic structure is designed for a practical application.

6.1 Distributed Couplings

Suppose that n propagating waves exist along a transmission system and other waves can be neglected, either because they are in the cut-off frequency ranges or because they are not excited. Let us indicate the magnitudes and phases of the n waves by complex numbers $a_1(z)$, $a_2(z)$, ..., $a_n(z)$ and let $\mathbf{a}(z)$ be a vector having the $a_n(z)$'s as its components. We assume that the waves are orthogonal to each other and normalized so that $p_i |a_i(z)|^2$ gives the transmission power of the i th wave in the positive z -direction, where p_i is 1 or -1 , and that the total transmission power is equal to the sum of individual transmission powers. In other words, the total power is given by

$$P = \sum p_i |a_i(z)|^2 = \mathbf{a}^+(z) \mathbf{P} \mathbf{a}(z) \quad (6.1)$$

where

$$\mathbf{P} = \text{diag} [p_1, p_2, \dots, p_n] \quad (6.2)$$

and $\mathbf{a}^+(z)$ is an abbreviation for $\{\mathbf{a}(z)\}^+$. In Chapter 3 the waves in a guide were shown to have this property.

Now, let us introduce some couplings between the n independent waves. The wall losses discussed in Section 3.8 certainly provide an example of coupling, however, in the present discussion we shall restrict ourselves to lossless couplings. An example of lossless coupling is shown in Fig. 6.1 in which two waveguides are coupled together through a large number of small windows in the common sidewall. If each waveguide has only one propagating mode, and if the waves propagating in the negative z -direction can be

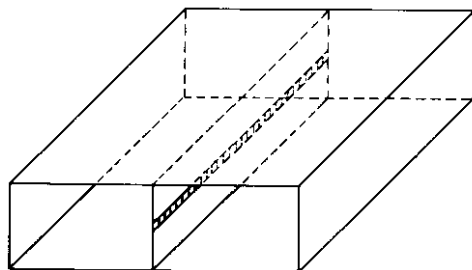


Fig. 6.1. An example of coupled circuits.

neglected because of the method of excitation, $a_1(z)$ may represent the wave propagating in the positive z -direction in one waveguide, and $a_2(z)$ in the other. The original unperturbed transmission lines in this case are two waveguides in parallel, each supporting only one propagating mode. Strictly speaking, a finite number of waves cannot satisfy the boundary conditions which are modified by couplings; however, for minor modifications, the effect of other waves on the transmission power may be negligible, and hence the total power is still given by (6.1) to a first order approximation. In the following discussion, this is assumed to be the case, and we shall study how $\mathbf{a}(z)$ changes with z .

The vector $\mathbf{a}(z)$ represents the magnitudes and phases of the waves at z , and hence $\mathbf{a}(z + \Delta z)$ represents those at $z + \Delta z$. Since the relation between the electromagnetic fields represented by $\mathbf{a}(z)$ and $\mathbf{a}(z + \Delta z)$ is linear, there must be a linear relation between $\mathbf{a}(z)$ and $\mathbf{a}(z + \Delta z)$. When higher order terms of Δz are neglected, this linear relation should be expressible in the form

$$\mathbf{a}(z + \Delta z) = \mathbf{a}(z) - \mathbf{C}\mathbf{a}(z) \Delta z \quad (6.3)$$

since $\mathbf{a}(z + \Delta z)$ should coincide with $\mathbf{a}(z)$ in the limit of $\Delta z \rightarrow 0$. The square matrix $\mathbf{C}$ thus introduced is called the coupling matrix. The negative sign in front of $\mathbf{C}$ is conveniently chosen so as to enable us to reach an eigenvalue problem in a conventional form. Transposing $\mathbf{a}(z)$ to the left-hand side, and taking the limit of $\Delta z \rightarrow 0$ after dividing both sides by Δz , we obtain a matrix differential equation

$$\frac{d\mathbf{a}(z)}{dz} = -\mathbf{C}\mathbf{a}(z) \quad (6.4)$$

The transmission power in the positive z -direction is given by (6.1), and this must be independent of z when the system is lossless. Therefore, we have

$$\frac{d}{dz} \{\mathbf{a}^+(z) \mathbf{P} \mathbf{a}(z)\} = 0 \quad (6.5)$$

Using the standard formula for differentiating a product of functions, we can write (6.5) as follows.

$$\frac{d\mathbf{a}^+(z)}{dz} \mathbf{P} \mathbf{a}(z) + \mathbf{a}^+(z) \mathbf{P} \frac{d\mathbf{a}(z)}{dz} = 0$$

Substituting (6.4) into this expression gives

$$-\mathbf{a}^+(z) (\mathbf{C}^+ \mathbf{P} + \mathbf{P} \mathbf{C}) \mathbf{a}(z) = 0 \quad (6.6)$$

Since this equation holds regardless of the value of $\mathbf{a}(z)$, we obtain

$$\mathbf{C}^+ \mathbf{P} + \mathbf{P} \mathbf{C} = 0 \quad (6.7)$$

This is the condition which the lossless coupling matrix $\mathbf{C}$ has to satisfy.

From (6.7), the diagonal components of $\mathbf{C}$ must satisfy

$$C_{ii} = -C_{ii}^* \quad (6.8)$$

in other words, the diagonal components of $\mathbf{C}$ are all pure imaginary. To evaluate the magnitude of C_{ii} , let us assume that the magnitudes of all the waves except $a_i(z)$ are zero at z , then from (6.4) we have

$$\frac{da_i(z)}{dz} = -C_{ii}a_i(z)$$

Suppose that $a_i(z)$ changes with z as $\exp(-j\beta_i z)$ before the couplings are introduced; the introduction of small couplings represented by $\mathbf{C}$ should, in general, cause $a_i(z)$ to behave somewhat differently. However, when no other waves exist at z , as we have assumed, no effects are expected through the couplings, and hence $a_i(z)$ is expected to change as $\exp(-j\beta_i z)$. A substitution of this functional form into the above equation shows that C_{ii} must be approximated by $j\beta_i$ which is pure imaginary as required form (6.8).

From (6.7) the off-diagonal components of $\mathbf{C}$ have to satisfy

$$C_{ij} = -C_{ji}^* \quad (p_i = p_j) \quad (6.9)$$

$$C_{ij} = C_{ji}^* \quad (p_i = -p_j) \quad (6.10)$$

The first condition applies when p_i and p_j have the same sign while the second one applies when the signs are opposite. Note that these conditions impose a restriction on the relation between C_{ij} and C_{ji} but no direct restrictions on the value of C_{ij} itself.

In order to derive an appropriate eigenvalue problem from (6.4), let us assume that $\mathbf{a}(z)$ changes exponentially with z , i.e., $\mathbf{a}(z) = \mathbf{x}e^{-\gamma z}$, as we did in the discussion of waveguides. Substituting this expression into (6.4), we obtain the desired equation

$$(\mathbf{C} - \gamma \mathbf{I}) \mathbf{x} = 0 \quad (6.11)$$

It is expected that any solution of (6.4) may be expressible as a linear combination of the solutions of the eigenvalue problem (6.11); the behavior of these solutions is simple and well understood, and hence a systematic eigenfunction approach to the coupling problem should become possible.

For this reason, we shall proceed with a study of the above eigenvalue problem and establish the basis for the eigenfunction approach.

First, we note that the discussion of the eigenvalue problem of self-adjoint matrices given in Section 5.1 is not applicable to the present problem, since C^+ is not equal to C . Using the relation (6.7), however, we can derive several important theorems. Equation (6.11) represents n simultaneous linear equations for the components of $\mathbf{x}$. To ensure nontrivial solutions, the determinant of the coefficients must equal zero which leads to an algebraic equation of the n th degree for γ whose roots give the eigenvalues,

$$\det(C - \gamma I) = 0 \quad (6.12)$$

Let us consider the case in which all the roots are distinct, then the following four theorems I-IV can be derived. The proof will be given after each statement.

I. Any n -dimensional vector can be expressed as a linear combination of the eigenvectors.

Proof. Since each root of (6.12) gives an eigenvector, we obtain n different eigenvectors. Suppose that the eigenvectors $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{p-1}$ are independent (i.e., $\sum_{i=1}^{p-1} \alpha_i \mathbf{x}_i = 0$ cannot be satisfied unless the α_i 's are all zero), but that $\mathbf{x}_p$ ($p \leq n$) is not, then we have

$$\mathbf{x}_p = \sum_{k=1}^{p-1} \alpha_k \mathbf{x}_k \quad (6.13)$$

Applying C from the left, we obtain

$$\gamma_p \mathbf{x}_p = \sum_{k=1}^{p-1} \alpha_k \gamma_k \mathbf{x}_k \quad (6.14)$$

Substitution of (6.13) into (6.14) gives

$$\sum_{k=1}^{p-1} (\gamma_k - \gamma_p) \alpha_k \mathbf{x}_k = 0 \quad (6.15)$$

By hypothesis, $\gamma_k \neq \gamma_p$ and not all of the α_k 's are equal to zero, (6.15) therefore shows that $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{p-1}$ are not independent of each other which leads to a contradiction. Thus, $\mathbf{x}_p$ is proved to be independent of $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_{p-1}$ provided that these $p-1$ eigenvectors are independent of each other. Proceeding from $p=2$ and increasing p in unit steps, we see that all the eigenvectors are independent of each other. Since an arbitrary vector can be expressed as a linear combination of n independent vectors in an n -dimen-

sional space, and since the eigenvectors form a set of n independent vectors, any vector can be expressed as a linear combination of the eigenvectors.

II. If γ_k is an eigenvalue, $-\gamma_k^*$ also becomes an eigenvalue.

Proof. Since γ_k is an eigenvalue, we have

$$\det(\mathbf{C} - \gamma_k \mathbf{I}) = 0$$

Taking the complex conjugate transposed determinant, we obtain

$$\det(\mathbf{C}^+ - \gamma_k^* \mathbf{I}) = 0$$

Multiplying $\det \mathbf{P}^{-1}$ and $\det \mathbf{P}$ from the left and from the right, respectively, the above equation becomes

$$\det(\mathbf{P}^{-1} \mathbf{C}^+ \mathbf{P} - \gamma_k^* \mathbf{I}) = 0$$

Replacing $\mathbf{C}^+ \mathbf{P}$ by $-\mathbf{PC}$, which are equal to each other from (6.7), we obtain

$$\det(\mathbf{C} + \gamma_k^* \mathbf{I}) = 0$$

which shows that $-\gamma_k^*$ is an eigenvalue of $\mathbf{C}$. When γ_k is pure imaginary, the theorem is trivial since γ_k is equal to $-\gamma_k^*$, however, when γ_k has a real part, the theorem asserts that an eigenvalue having the opposite sign in front of the real part also exists. In other words, whenever a mode exists which grows with distance, a corresponding decaying mode exists having the same phase constant.

III. Let γ_k and γ_l be the eigenvalues corresponding to eigenvectors $\mathbf{x}_k$ and $\mathbf{x}_l$, respectively. If $\gamma_k \neq -\gamma_l^*$, there is an orthogonal relation between $\mathbf{x}_l$ and $\mathbf{x}_k$ of the form

$$\mathbf{x}_l^+ \mathbf{P} \mathbf{x}_k = 0 \quad (6.16)$$

Proof. By definition, we have

$$\mathbf{C} \mathbf{x}_k = \gamma_k \mathbf{x}_k, \quad \mathbf{C} \mathbf{x}_l = \gamma_l \mathbf{x}_l$$

Multiplying the first equation by $\mathbf{x}_l^+ \mathbf{P}$ from the left, and the adjoint of the second equation by $\mathbf{P} \mathbf{x}_k$ from the right, and then adding the two equations, we have

$$\mathbf{x}_l^+ \mathbf{P} \mathbf{C} \mathbf{x}_k + \mathbf{x}_l^+ \mathbf{C}^+ \mathbf{P} \mathbf{x}_k = (\gamma_k + \gamma_l^*) \mathbf{x}_l^+ \mathbf{P} \mathbf{x}_k$$

The left-hand side can be rewritten in the form $\mathbf{x}_l^+ (\mathbf{PC} + \mathbf{C}^+ \mathbf{P}) \mathbf{x}_k$, which is equal to zero from (6.7). Therefore, the right-hand side must also vanish which means that (6.16) holds, since $\gamma_k + \gamma_l^*$ is not equal to zero by hypothesis.

IV. If $\gamma_k = -\gamma_l^*$, we have

$$\mathbf{x}_l^+ \mathbf{P} \mathbf{x}_k \neq 0 \quad (6.17)$$

Proof. By theorem I, $\mathbf{P} \mathbf{x}_l$ can be expressed as a linear combination of the eigenvectors, i.e.,

$$\mathbf{P} \mathbf{x}_l = \sum_m \alpha_m \mathbf{x}_m$$

where the summation is extended from $m=1$ to n . Multiplying the above equation by $(\mathbf{P} \mathbf{x}_l)^+ = \mathbf{x}_l^+ \mathbf{P}$ from the left, and since $\mathbf{x}_l$ is not equal to zero, we have

$$0 \neq (\mathbf{P} \mathbf{x}_l)^+ (\mathbf{P} \mathbf{x}_l) = \sum_m \alpha_m \mathbf{x}_l^+ \mathbf{P} \mathbf{x}_m$$

However, all the terms on the right-hand side, except the k th term for which $\gamma_k = -\gamma_l^*$, are equal to zero. Therefore, we obtain

$$\alpha_k \mathbf{x}_l^+ \mathbf{P} \mathbf{x}_k \neq 0$$

from which (6.17) follows immediately.

When γ_k is pure imaginary, let us define $\tilde{\mathbf{x}}_k$ to be

$$\tilde{\mathbf{x}}_k = c \mathbf{x}_k^+ \quad (6.18)$$

On the other hand, when γ_k has a nonzero real part

$$\tilde{\mathbf{x}}_k = c \mathbf{x}_l^+ \quad (6.19)$$

where $\mathbf{x}_l$ is the eigenvector corresponding to $\gamma_l = -\gamma_k^*$ whose existence is guaranteed by Theorem II. It follows from Theorem IV that $\tilde{\mathbf{x}}_k \mathbf{P} \mathbf{x}_k$ is not equal to zero provided the constant c is not equal to zero. Let us choose the value of c in such a way that $\tilde{\mathbf{x}}_k \mathbf{P} \mathbf{x}_k$ becomes unity, then we have the orthogonality and normalization conditions

$$\tilde{\mathbf{x}}_l \mathbf{P} \mathbf{x}_k = 0 \quad (l \neq k) \quad (6.20)$$

$$\tilde{\mathbf{x}}_k \mathbf{P} \mathbf{x}_k = 1 \quad (6.21)$$

The normalization condition does not uniquely determine $\mathbf{x}_k$; if one wants to determine it uniquely, another condition must be imposed such as $\mathbf{x}_k^+ \mathbf{x}_k = 1$, in our present discussion, however, this is not necessary.

The above results are derived under the assumption that all the n eigenvalues are distinct, however, (6.12) may have one or more multiple roots. In such a case, the number of distinct eigenvalues is less than n , and the above technique of proving Theorem I fails. Fortunately, we can still obtain n

independent eigenvectors, as we shall discuss next, and hence the theorem itself holds without modification.

Let $\mathbf{C}'$ be a coupling matrix satisfying (6.7), and let us consider $\mathbf{C} + \varepsilon\mathbf{C}'$, where $\mathbf{C}$ is the coupling matrix under consideration, and ε is a small positive number. The n eigenvalues of $\mathbf{C} + \varepsilon\mathbf{C}'$ can be made distinct from one another by choosing $\mathbf{C}'$ and ε properly; the eigenvectors $\mathbf{x}_k(\varepsilon)$ of $\mathbf{C} + \varepsilon\mathbf{C}'$ and the corresponding $\tilde{\mathbf{x}}_k(\varepsilon)$'s will then satisfy (6.20) and (6.21). The eigenvalue $\gamma_k(\varepsilon)$, corresponding to $\mathbf{x}_k(\varepsilon)$, is a root of

$$\det(\mathbf{C} + \varepsilon\mathbf{C}' - \gamma\mathbf{I}) = 0$$

The coefficients of this algebraic equation vary continuously with ε , hence the root also changes continuously with ε since it is a continuous function of the coefficients. In other words, $\gamma_k(\varepsilon)$ is a continuous function ε .

The components of $\mathbf{x}_k(\varepsilon)$ are given by the solutions of n simultaneous equations

$$(\mathbf{C} + \varepsilon\mathbf{C}' - \gamma_k(\varepsilon)\mathbf{I})\mathbf{x} = 0$$

Since $\gamma_k(\varepsilon)$ is assumed to be a simple root, the ratios of the components of $\mathbf{x}_k(\varepsilon)$ are given by

$$-D_1 : -D_2 : \cdots : -D_{n-1} : D$$

where D is the determinant of the coefficients of the above simultaneous equations with its n th row and n th column removed, and D_i is this modified determinant except that the i th column is made up of the last coefficient in each equation. Therefore, the ratios also vary continuously with ε .

If we define γ_k , $\mathbf{x}_k$, and $\tilde{\mathbf{x}}_k$ by

$$\gamma_k = \lim_{\varepsilon \rightarrow 0} \gamma_k(\varepsilon), \quad \mathbf{x}_k = \lim_{\varepsilon \rightarrow 0} \mathbf{x}_k(\varepsilon), \quad \tilde{\mathbf{x}}_k = \lim_{\varepsilon \rightarrow 0} \tilde{\mathbf{x}}_k(\varepsilon)$$

respectively, then γ_k , the component ratios of $\mathbf{x}_k$ and those of $\tilde{\mathbf{x}}_k$ are uniquely determined since they are all continuous functions of ε . Furthermore, γ_k and $\mathbf{x}_k$ respectively become an eigenvalue of (6.11) and the corresponding eigenvector. Since the $\mathbf{x}_k(\varepsilon)$'s and the $\tilde{\mathbf{x}}_k(\varepsilon)$'s satisfy (6.20) and (6.21) regardless of the value of ε , the $\mathbf{x}_k$'s and the $\tilde{\mathbf{x}}_k$'s thus defined also satisfy these two equations. The $\mathbf{x}_k$'s are independent of each other from the proof of theorem I provided the corresponding eigenvalues are distinct. However, for the degenerate eigenvectors, we have to modify the discussion as follows. Suppose that they are not independent, then we have

$$\sum \alpha_i \mathbf{x}_i = 0$$

where the $\mathbf{x}_i$'s are degenerate eigenvectors and the α_i 's are not all equal to zero. Assuming α_k is not equal to zero, and multiplying by $\tilde{\mathbf{x}}_k \mathbf{P}$ from the left, we have

$$\alpha_k \tilde{\mathbf{x}}_k \mathbf{P} \mathbf{x}_k = 0$$

which contradicts (6.21). Thus, we see that the n eigenvectors are independent of each other, and Theorems I-IV hold without modification, even in degenerate cases.

It may be worth noting that different sets of eigenvectors will be obtained depending on the choice of $\mathbf{C}'$ in the above procedure. However, each set satisfies (6.20) and (6.21), and there is no preference for any particular set, unless we wish to introduce another perturbation in the system.

An arbitrary vector can be expressed as a linear combination of the n independent eigenvectors:

$$\mathbf{x} = \sum \alpha_k \mathbf{x}_k$$

If we multiply this equation by $\tilde{\mathbf{x}}_k \mathbf{P}$ from the left, all the terms on the right-hand side, except the k th one, disappear because of (6.20) and (6.21), leaving

$$\alpha_k = \tilde{\mathbf{x}}_k \mathbf{P} \mathbf{x}$$

A substitution into the original expression for $\mathbf{x}$ gives

$$\mathbf{x} = \sum \mathbf{x}_k (\tilde{\mathbf{x}}_k \mathbf{P} \mathbf{x}) \quad (6.22)$$

Multiplying by $\mathbf{C}$ from the left and using $\mathbf{C} \mathbf{x}_k = \gamma_k \mathbf{x}_k$, we have

$$\mathbf{C} \mathbf{x} = \sum \gamma_k \mathbf{x}_k (\tilde{\mathbf{x}}_k \mathbf{P} \mathbf{x}) \quad (6.23)$$

Since (6.23) holds regardless of the value of $\mathbf{x}$, $\mathbf{C}$ can be written in the form

$$\mathbf{C} \cdot = \sum \gamma_k \mathbf{x}_k \tilde{\mathbf{x}}_k \mathbf{P} \cdot \quad (6.24)$$

which is called the spectral representation of $\mathbf{C}$, where the dot signifies that $\mathbf{C}$ is an operator.

6.2 Perturbation Method for the Eigenvalue Problem

To obtain the solutions of (6.11), we have to solve an algebraic equation of n th degree. In principle, this is always possible, but, in practice, it may prove to be a formidable problem. Fortunately, when the couplings are small, good approximate solutions can be obtained by a perturbation method. In this section, we shall study this method in detail, assuming that the couplings are small. Two distinct types of wave interaction will be shown to

exist depending on the signs of the transmission powers of interacting waves.

Let $\mathbf{C}_0$ be a diagonal matrix having diagonal components identical to those of $\mathbf{C}$ itself, and let $\mathbf{C}_1$ be $\mathbf{C} - \mathbf{C}_0$, then the diagonal components of $\mathbf{C}_1$ are all zero, and its off-diagonal components are small. We write $\mathbf{C}_0$ and $\mathbf{C}_1$ as follows:

$$\mathbf{C}_0 = \begin{bmatrix} C_{11} & 0 & \dots & 0 \\ 0 & C_{22} & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & C_{nn} \end{bmatrix}, \quad \mathbf{C}_1 = \begin{bmatrix} 0 & C_{12} & \dots & C_{1n} \\ C_{21} & 0 & \dots & C_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ C_{n1} & \dots & \dots & 0 \end{bmatrix} \quad (6.25)$$

By inspection, the eigenvalues of $\mathbf{C}_0$ are found to be $C_{11}, C_{22}, \dots, C_{nn}$, and the eigenvectors are given by

$$\mathbf{x}_1^{(0)} = \begin{bmatrix} 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, \quad \mathbf{x}_2^{(0)} = \begin{bmatrix} 0 \\ 1 \\ \vdots \\ 0 \end{bmatrix}, \quad \dots, \quad \mathbf{x}_n^{(0)} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \end{bmatrix} \quad (6.26)$$

The corresponding $\mathbf{x}_i^{(0)}$'s are

$$\mathbf{x}_1^{(0)} = [p_1 \ 0 \ \dots \ 0], \quad \mathbf{x}_2^{(0)} = [0 \ p_2 \ \dots \ 0], \dots \quad (6.27)$$

Since $\mathbf{C}_1$ is small, there may be an eigenvector of $\mathbf{C}$, namely $\mathbf{x}_i$, which is similar to $\mathbf{x}_i^{(0)}$. Let us write $\mathbf{x}_i$ in the form

$$\mathbf{x}_i = \mathbf{x}_i^{(0)} + \sum_{n \neq i} c_n \mathbf{x}_n^{(0)} \quad (6.28)$$

The eigenvalue γ_i is expected to be close to $\gamma_i^{(0)} = C_{ii}$, hence we write γ_i in the form

$$\gamma_i = \gamma_i^{(0)} + \Delta\gamma_i \quad (6.29)$$

Substituting (6.28) and (6.29) into the left-hand side of (6.11), we obtain

$$(\mathbf{C} - \gamma_i \mathbf{I}) \mathbf{x}_i \simeq (\mathbf{C}_1 - \Delta\gamma_i) \mathbf{x}_i^{(0)} + \sum_{n \neq i} c_n (\gamma_n^{(0)} - \gamma_i^{(0)}) \mathbf{x}_n^{(0)} \quad (6.30)$$

where all the product terms between $\mathbf{C}_1$, $\Delta\gamma_i$, and c_n are neglected since they are all assumed to be small. Since $\mathbf{x}_i$ is an eigenvector of $\mathbf{C}$ with eigenvalue γ_i , the left-hand side is equal to zero. Multiplying (6.30) by $\mathbf{x}_i^{(0)\dagger}$ from the left and using the orthonormal property of the $\mathbf{x}_n^{(0)}$'s, we have

$$\mathbf{x}_i^{(0)\dagger} \mathbf{C}_1 \mathbf{x}_i^{(0)} \simeq \Delta\gamma_i \quad (6.31)$$

The left-hand side is equal to zero because the diagonal components of $\mathbf{C}_1$ are all zero. To a first order approximation, therefore, γ_i is equal to the original eigenvalue $\gamma_i^{(0)}$. Next, multiplying (6.30) by $\tilde{\mathbf{x}}_n^{(0)}\mathbf{P}$ from the left and using the orthonormal property, we obtain

$$c_n \simeq -\frac{\tilde{\mathbf{x}}_n^{(0)}\mathbf{P}\mathbf{C}_1\mathbf{x}_i^{(0)}}{\gamma_n^{(0)} - \gamma_i^{(0)}} \quad (n \neq i) \quad (6.32)$$

Substituting this into (6.28) gives the first order approximation of $\mathbf{x}_i$.

If one of the $\gamma_n^{(0)}$'s, say $\gamma_j^{(0)}$ ($j \neq i$), is equal to or close to $\gamma_i^{(0)}$, the right-hand side of (6.32) for $n=j$ becomes infinite or large, and the assumption that the c_n 's are small is no longer valid. In such a case, we proceed as in the discussion of the waveguide wall losses, and write $\mathbf{x}_i$ in the form

$$\mathbf{x}_i = A\mathbf{x}_i^{(0)} + B\mathbf{x}_j^{(0)} + \sum_{n \neq i, j} c_n \mathbf{x}_n^{(0)} \quad (6.33)$$

where A , B , and the c_n 's are constants to be determined. Let $2\Delta\gamma_0$ be $\gamma_i^{(0)} - \gamma_j^{(0)}$ and $2\gamma_0$ be $\gamma_i^{(0)} + \gamma_j^{(0)}$, then

$$\gamma_i^{(0)} = \gamma_0 + \Delta\gamma_0, \quad \gamma_j^{(0)} = \gamma_0 - \Delta\gamma_0 \quad (6.34)$$

Assuming that γ_i is close to γ_0 , we write it in the form

$$\gamma_i = \gamma_0 + \Delta\gamma_i \quad (6.35)$$

and determine $\Delta\gamma_i$. To do so, we first substitute (6.33), (6.34), and (6.35) into the left-hand side of (6.11) and obtain

$$\begin{aligned} (\mathbf{C} - \gamma_i\mathbf{I})\mathbf{x}_i \simeq & \mathbf{C}_1 \{A\mathbf{x}_i^{(0)} + B\mathbf{x}_j^{(0)}\} - (\Delta\gamma_i - \Delta\gamma_0)A\mathbf{x}_i^{(0)} \\ & - (\Delta\gamma_i + \Delta\gamma_0)B\mathbf{x}_j^{(0)} + \sum_{n \neq i, j} c_n (\gamma_n^{(0)} - \gamma_0) \mathbf{x}_n^{(0)} \end{aligned} \quad (6.36)$$

where the products of small terms are again neglected. The left-hand side is equal to zero. Multiplying (6.36) from the left by $\tilde{\mathbf{x}}_i^{(0)}\mathbf{P}$ and by $\tilde{\mathbf{x}}_j^{(0)}\mathbf{P}$ in turn, we obtain

$$-A(\Delta\gamma_i - \Delta\gamma_0) + B\tilde{\mathbf{x}}_i^{(0)}\mathbf{P}\mathbf{C}_1\mathbf{x}_j^{(0)} = 0 \quad (6.37)$$

$$A\tilde{\mathbf{x}}_j^{(0)}\mathbf{P}\mathbf{C}_1\mathbf{x}_i^{(0)} - B(\Delta\gamma_i + \Delta\gamma_0) = 0 \quad (6.38)$$

where the orthonormal property between the $\mathbf{x}_n^{(0)}$'s as well as $\tilde{\mathbf{x}}_i^{(0)}\mathbf{P}\mathbf{C}_1\mathbf{x}_i^{(0)} = 0$ and $\tilde{\mathbf{x}}_j^{(0)}\mathbf{P}\mathbf{C}_1\mathbf{x}_j^{(0)} = 0$ is used. If the simultaneous equations (6.37) and (6.38) for A and B are to have nontrivial solutions, the determinant of the coefficients must vanish. This condition gives

$$\Delta\gamma_i^2 = \tilde{\mathbf{x}}_i^{(0)}\mathbf{P}\mathbf{C}_1\mathbf{x}_j^{(0)} \cdot \tilde{\mathbf{x}}_j^{(0)}\mathbf{P}\mathbf{C}_1\mathbf{x}_i^{(0)} + \Delta\gamma_0^2$$

and $\tilde{\mathbf{x}}_i^{(0)} \mathbf{P} \mathbf{C}_1 \mathbf{x}_j^{(0)}$ is calculated to be C_{ij} and $\tilde{\mathbf{x}}_j^{(0)} \mathbf{P} \mathbf{C}_1 \mathbf{x}_i^{(0)}$ is equal to C_{ji} , hence the first term on the right-hand side becomes $C_{ij} C_{ji}$, which is equal to $-|C_{ij}|^2$ or $|C_{ij}|^2$ from (6.9) and (6.10) depending on the sign of $p_i p_j$. Since $\Delta\gamma_0$ is always pure imaginary, let us write it as $j\Delta\beta_0$ where $\Delta\beta_0$ is real, then the second term on the right-hand side becomes $-(\Delta\beta_0)^2$. Thus, we have

$$\Delta\gamma_i^2 = -|C_{ij}|^2 - (\Delta\beta_0)^2 \quad (p_i = p_j) \quad (6.39)$$

$$\Delta\gamma_i^2 = |C_{ij}|^2 - (\Delta\beta_0)^2 \quad (p_i = -p_j) \quad (6.40)$$

When $p_i = p_j$, $\Delta\gamma_i$, and hence γ_i become pure imaginary. On the other hand, when $p_i = -p_j$ and $|C_{ij}|^2 > (\Delta\beta_0)^2$, $\Delta\gamma_i$ becomes real, and hence γ_i becomes complex. Since (6.40) only determines the square of $\Delta\gamma_i$, two different γ_i 's with opposite signs in front of the real part are obtained, as Theorem II in Section 6.1 indicated. Figures 6.2(a) and (b) illustrate the relation between $\gamma_i = \alpha + j\beta$ and $\Delta\beta_0$ for the cases in which $p_i = p_j$ and $p_i = -p_j$, respectively, where α and β are the real and imaginary parts of γ_i . The term β_0 is the mean value while $\Delta\beta_0$ is the difference of the phase constants of the two original unperturbed waves.

The ratio of A to B is calculated from (6.37) to be

$$\frac{A}{B} = \frac{\tilde{\mathbf{x}}_i^{(0)} \mathbf{P} \mathbf{C}_1 \mathbf{x}_j^{(0)}}{\Delta\gamma_i - \Delta\gamma_0} \quad (6.41)$$

Corresponding to two different values of $\Delta\gamma_i$, two different ratios are obtained. The value of c_n is easily determined by multiplying (6.36) by $\tilde{\mathbf{x}}_n^{(0)} \mathbf{P}$ from the left:

$$c_n = - \frac{\tilde{\mathbf{x}}_n^{(0)} \mathbf{P} \mathbf{C}_1 \{A \mathbf{x}_i^{(0)} + B \mathbf{x}_j^{(0)}\}}{\gamma_n^{(0)} - \gamma_0} \quad (6.42)$$

When the $\gamma_n^{(0)}$'s ($n \neq i, j$) are not close to γ_0 , the c_n 's are all small and can be neglected since they only introduce second order variations to the transmission powers. On the other hand, $\Delta\gamma_i$ cannot be neglected, even though its magnitude is small since its effect appears as the exponential of $\Delta\gamma_i z$ which becomes large as z increases. The wave amplitude corresponding to real and negative $\Delta\gamma_i$ grows with z . This suggests that there is a possibility of obtaining power amplification when a small coupling is introduced between two waves propagating in the same direction with nearly equal propagation constants but with positive and negative transmission powers. In ordinary

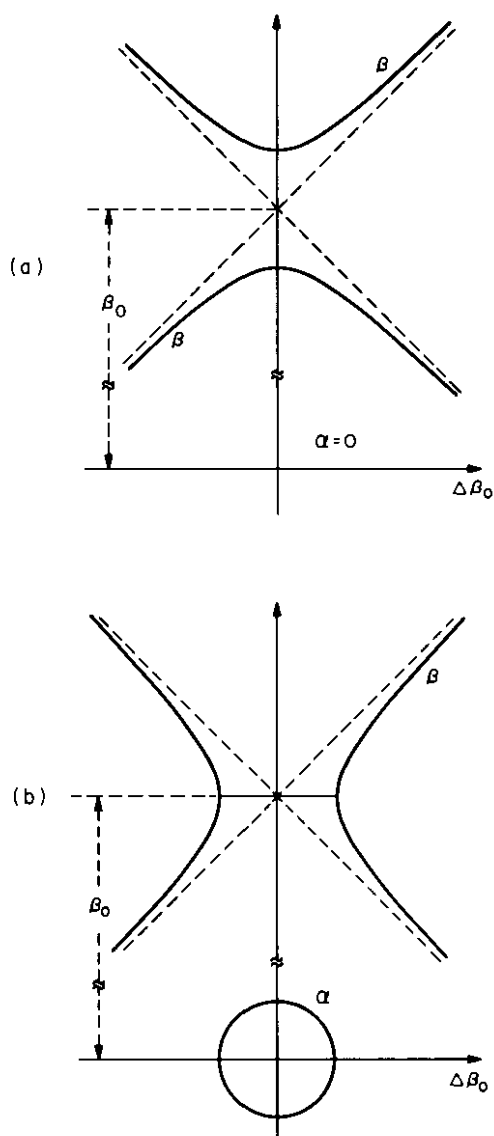


Fig. 6.2. The components α and β of the propagation constant γ as functions of $\Delta\beta_0$.

passive networks, no such amplification takes place since waves do not carry negative transmission powers. An electron beam, however, can support waves with negative transmission powers, and amplification does take place when one such wave is coupled to a wave propagating along a nearby passive circuit with a similar propagation constant as we shall discuss in Chapter 8.

From (6.31) and (6.32), we saw that waves with sufficiently different propagation constants do not interact with each other under the small coupling condition. This is understandable since the phase relation between waves changes rapidly with distance, and if there is some kind of interaction at one point, just the opposite kind of interaction takes place at another point thus causing the effects to cancel each other. On the other hand, if the propagation constants are similar, the effects of interaction accumulate with distance to such an extent that the originally independent waves can no longer be considered as independent, and a pair of new waves with different propagation constants result.

In the above discussion, only two waves were assumed to have similar propagation constants, however, a similar method can be applied to cases in which the propagation constants of three or more waves are similar. In such cases, algebraic equations of correspondingly higher degree have to be solved. An example of three-wave interaction will be given in Chapter 8 in connection with coupling between an electron beam and a transmission circuit.

6.3 Interactions between Two Waves

As a simple, but important, example of distributed coupling, let us consider the interactions between two waves propagating in the positive z -direction. When there are several coupled waves, only waves with similar propagation constants interact, and the remainder can be considered to propagate independently, as we learned in the previous section. Since there are a number of cases in which only two waves have similar propagation constants, our present study has definite practical value besides having the advantage of mathematical simplicity.

Let $a_1(z)$ represent one wave at z and $a_2(z)$ the other, and let $\beta_0 + \Delta\beta_0$ and $\beta_0 - \Delta\beta_0$ be their respective propagation constants. First, we shall consider the case in which the total transmission power in the positive z -direction is given by $|a_1(z)|^2 + |a_2(z)|^2$. In this case, $p_1 = p_2 = 1$ and the coupling coefficients C_{12} and C_{21} satisfy the relation $C_{12} = -C_{21}^*$ from

(6.9). Thus we have $\mathbf{P} = \text{diag}[1 \ 1]$, and

$$\mathbf{C}_0 = \begin{bmatrix} j(\beta_0 + \Delta\beta_0) & 0 \\ 0 & j(\beta_0 - \Delta\beta_0) \end{bmatrix}, \quad \mathbf{C}_1 = \begin{bmatrix} 0 & C_{12} \\ -C_{12}^* & 0 \end{bmatrix}$$

The eigenvectors of $\mathbf{C}_0$ are given by

$$\mathbf{x}_1^{(0)} = \begin{bmatrix} 1 \\ 0 \end{bmatrix}, \quad \mathbf{x}_2^{(0)} = \begin{bmatrix} 0 \\ 1 \end{bmatrix}$$

and the corresponding $\tilde{\mathbf{x}}_1^{(0)}$ and $\tilde{\mathbf{x}}_2^{(0)}$ are given by

$$\tilde{\mathbf{x}}_1^{(0)} = [1 \ 0], \quad \tilde{\mathbf{x}}_2^{(0)} = [0 \ 1]$$

Let us write $\Delta\gamma_i = \pm j \Delta\beta$, then we have from (6.39)

$$\Delta\beta = \{|C_{12}|^2 + (\Delta\beta_0)^2\}^{1/2} \quad (6.43)$$

The eigenvalues of $\mathbf{C}_0 + \mathbf{C}_1$ are approximately given by $j(\beta_0 + \Delta\beta)$ and $j(\beta_0 - \Delta\beta)$, and the corresponding eigenvectors are calculated to be

$$\mathbf{x}_1 = \begin{bmatrix} 1 \\ j(\Delta\beta - \Delta\beta_0)/C_{12} \end{bmatrix}, \quad \mathbf{x}_2 = \begin{bmatrix} 1 \\ -j(\Delta\beta + \Delta\beta_0)/C_{12} \end{bmatrix}$$

where use is made of (6.41). The corresponding $\tilde{\mathbf{x}}_1$ and $\tilde{\mathbf{x}}_2$ are

$$\tilde{\mathbf{x}}_1 = \begin{bmatrix} \frac{|C_{12}|^2}{|C_{12}|^2 + (\Delta\beta - \Delta\beta_0)^2} & \frac{-jC_{12}(\Delta\beta - \Delta\beta_0)}{|C_{12}|^2 + (\Delta\beta - \Delta\beta_0)^2} \end{bmatrix}$$

$$\tilde{\mathbf{x}}_2 = \begin{bmatrix} \frac{|C_{12}|^2}{|C_{12}|^2 + (\Delta\beta + \Delta\beta_0)^2} & \frac{jC_{12}(\Delta\beta + \Delta\beta_0)}{|C_{12}|^2 + (\Delta\beta + \Delta\beta_0)^2} \end{bmatrix}$$

Since we have obtained all the necessary quantities in their explicit forms, let us next study how $a_1(z)$ and $a_2(z)$ change with z , assuming that $a_1(z) = A_0$ and $a_2(z) = 0$ at $z = 0$. The condition at $z = 0$ can be expressed by

$$\mathbf{a}(0) = \begin{bmatrix} A_0 \\ 0 \end{bmatrix}$$

In terms of the eigenvectors $\mathbf{x}_1$ and $\mathbf{x}_2$, $\mathbf{a}(0)$ becomes

$$\begin{aligned} \mathbf{a}(0) &= \tilde{\mathbf{x}}_1 \mathbf{P} \begin{bmatrix} A_0 \\ 0 \end{bmatrix} \mathbf{x}_1 + \tilde{\mathbf{x}}_2 \mathbf{P} \begin{bmatrix} A_0 \\ 0 \end{bmatrix} \mathbf{x}_2 \\ &= A_0 |C_{12}|^2 \left\{ \frac{\mathbf{x}_1}{|C_{12}|^2 + (\Delta\beta - \Delta\beta_0)^2} + \frac{\mathbf{x}_2}{|C_{12}|^2 + (\Delta\beta + \Delta\beta_0)^2} \right\} \end{aligned}$$

The wave corresponding to $\mathbf{x}_1$ varies with z as $\exp\{-j(\beta_0 + \Delta\beta)z\}$ and that corresponding to $\mathbf{x}_2$ varies as $\exp\{-j(\beta_0 - \Delta\beta)z\}$. As a result, $\mathbf{a}(z)$ at arbitrary z becomes

$$\mathbf{a}(z) = \begin{bmatrix} a_1(z) \\ a_2(z) \end{bmatrix} = A_0 |C_{12}|^2 e^{-j\beta_0 z} \times \left\{ \frac{e^{-j\Delta\beta z}}{|C_{12}|^2 + (\Delta\beta - \Delta\beta_0)^2} \begin{bmatrix} 1 \\ j(\Delta\beta - \Delta\beta_0)/C_{12} \end{bmatrix} + \frac{e^{j\Delta\beta z}}{|C_{12}|^2 + (\Delta\beta + \Delta\beta_0)^2} \begin{bmatrix} 1 \\ -j(\Delta\beta + \Delta\beta_0)/C_{12} \end{bmatrix} \right\}$$

Therefore, $a_1(z)$ is given by

$$A_0 |C_{12}|^2 e^{-j\beta_0 z} \times \frac{|C_{12}|^2 (e^{-j\Delta\beta z} + e^{j\Delta\beta z}) + (\Delta\beta + \Delta\beta_0)^2 e^{-j\Delta\beta z} + (\Delta\beta - \Delta\beta_0)^2 e^{j\Delta\beta z}}{\{|C_{12}|^2 + (\Delta\beta - \Delta\beta_0)^2\} \{|C_{12}|^2 + (\Delta\beta + \Delta\beta_0)^2\}}$$

Using (6.43) the denominator and numerator can both be simplified to $4\Delta\beta^2 |C_{12}|^2$ and $4\Delta\beta^2 \cos \Delta\beta z - 4j\Delta\beta \Delta\beta_0 \sin \Delta\beta z$, respectively, and $a_1(z)$ becomes

$$a_1(z) = A_0 e^{-j\beta_0 z} \{\cos \Delta\beta z - j(\Delta\beta_0/\Delta\beta) \sin \Delta\beta z\}$$

To interpret this expression, we multiply by $\sqrt{2}e^{j\omega t}$ and take the real part:

$$\begin{aligned} \text{Re}\{\sqrt{2}a_1(z)e^{j\omega t}\} \\ = \sqrt{2}A_0 \{1 - (|C_{12}|/\Delta\beta)^2 \sin^2 \Delta\beta z\}^{1/2} \cos(\omega t - \beta_0 z - \varphi_1) \end{aligned} \quad (6.44)$$

where φ_1 is defined through

$$\tan \varphi_1 = (\Delta\beta_0/\Delta\beta) \tan \Delta\beta z \quad (6.45)$$

Similarly, for $a_2(z)$ we obtain

$$\text{Re}\{\sqrt{2}a_2(z)e^{j\omega t}\} = \sqrt{2}A_0 (|C_{12}|/\Delta\beta) \sin \Delta\beta z \cos(\omega t - \beta_0 z - \varphi_2) \quad (6.46)$$

where φ_2 is defined through

$$C_{12}^* = |C_{12}| e^{-j\varphi_2} \quad (6.47)$$

We see from (6.44) and (6.46) that both waves $a_1(z)$ and $a_2(z)$ propagate with the same phase constant β_0 in the positive z -direction, but their magnitudes fluctuate with z . Since the power carried by each wave is given by the square of the magnitude, it also fluctuates with z as shown in Fig. 6.3. When $\Delta\beta_0 = 0$, $|C_{12}|^2/\Delta\beta^2 = 1$, $|a_1(z)|^2 = 0$, and $|a_2(z)|^2 = A_0^2$ at $z =$

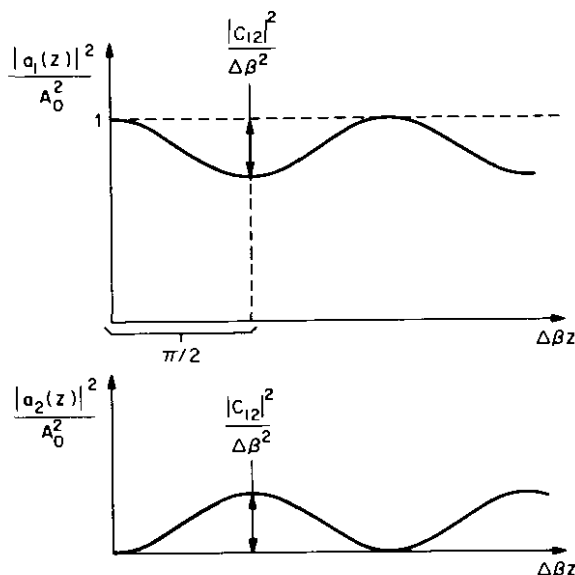


Fig. 6.3. Exchange of transmission power between two coupled modes.

$\{(2n-1)\pi/2 \Delta\beta\}$. Furthermore, at $z = (n\pi/\Delta\beta)$, $|a_1(z)|^2$ returns to A_0^2 and $|a_2(z)|^2$ becomes zero, where n is an arbitrary integer. In other words, when the original unperturbed waves have the same phase velocity, the power introduced into one wave at $z=0$ is completely transferred to the other wave and back again by the time it reaches $z = (\pi/2 \Delta\beta)$ and $z = (\pi/\Delta\beta)$, respectively. This exchange of power continues back and forth between the two waves until the coupling is removed at some point along the transmission system.

Let us next consider the case where the total transmission power is given by $|a_1(z)|^2 + |a_2(z)|^2$. In this case, $\mathbf{P} = \text{diag}[1 \ 1]$, and

$$\mathbf{C}_0 = \begin{bmatrix} j(\beta_0 + \Delta\beta_0) & 0 \\ 0 & j(\beta_0 - \Delta\beta_0) \end{bmatrix}, \quad \mathbf{C}_1 = \begin{bmatrix} 0 & C_{12} \\ C_{12}^* & 0 \end{bmatrix}$$

Suppose that $|C_{12}|^2 \geq (\Delta\beta_0)^2$, then from (6.40), we have

$$\alpha = \{|C_{12}|^2 - (\Delta\beta_0)^2\}^{1/2} \quad (6.48)$$

where $\Delta\gamma_i$ is indicated by $\pm\alpha$. The eigenvalues of $\mathbf{C}_0 + \mathbf{C}_1$ are given by

$\alpha + j\beta_0$ and $-\alpha + j\beta_0$. The corresponding eigenvectors are calculated to be

$$\mathbf{x}_1 = \begin{bmatrix} 1 \\ (\alpha - j\Delta\beta_0)/C_{12} \end{bmatrix}, \quad \mathbf{x}_2 = \begin{bmatrix} 1 \\ (-\alpha - j\Delta\beta_0)/C_{12} \end{bmatrix}$$

and $\tilde{\mathbf{x}}_1$ and $\tilde{\mathbf{x}}_2$ are given by

$$\tilde{\mathbf{x}}_1 = \begin{bmatrix} \frac{|C_{12}|^2}{|C_{12}|^2 + (-\alpha + j\Delta\beta_0)^2} & \frac{C_{12}(-\alpha + j\Delta\beta_0)}{|C_{12}|^2 + (-\alpha + j\Delta\beta_0)^2} \\ \frac{|C_{12}|^2}{|C_{12}|^2 + (\alpha + j\Delta\beta_0)^2} & \frac{C_{12}(\alpha + j\Delta\beta_0)}{|C_{12}|^2 + (\alpha + j\Delta\beta_0)^2} \end{bmatrix}$$

If $a_1(z) = A_0$ and $a_2(z) = 0$ at $z = 0$, as in the previous example, we have

$$\begin{bmatrix} a_1(z) \\ a_2(z) \end{bmatrix} = A_0 |C_{12}|^2 e^{-j\beta_0 z} \times \left\{ \frac{e^{-\alpha z}}{|C_{12}|^2 + (-\alpha + j\Delta\beta_0)^2} \mathbf{x}_1 + \frac{e^{\alpha z}}{|C_{12}|^2 + (\alpha + j\Delta\beta_0)^2} \mathbf{x}_2 \right\}$$

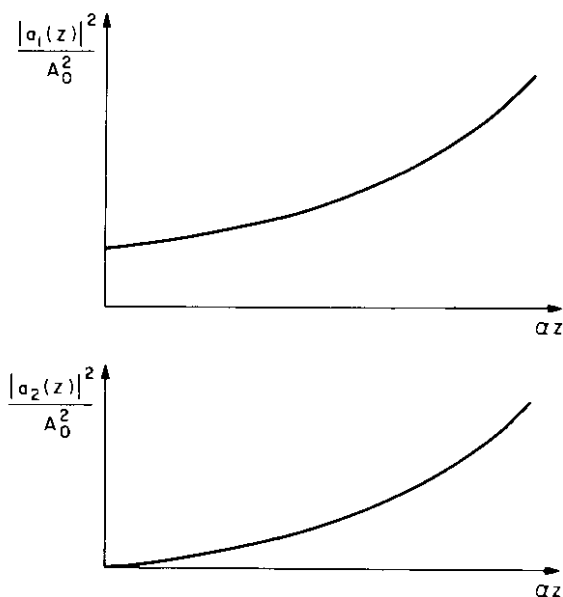


Fig. 6.4. Growing transmission power. Note that $|a_1(z)|^2 - |a_2(z)|^2$ is kept constant.

To interpret this result, we take the real parts of $\sqrt{2}e^{j\omega t}$ times $a_1(z)$ and $a_2(z)$:

$$\operatorname{Re} \{ \sqrt{2} a_1(z) e^{j\omega t} \} = \sqrt{2} A_0 \{ 1 + (|C_{12}|/\alpha)^2 \sinh^2 \alpha z \}^{1/2} \cos(\omega t - \beta_0 z - \varphi_1) \quad (6.49)$$

$$\operatorname{Re} \{ \sqrt{2} a_2(z) e^{j\omega t} \} = -\sqrt{2} A_0 (|C_{12}|/\alpha) \sinh \alpha z \cos(\omega t - \beta_0 z - \varphi_2) \quad (6.50)$$

where φ_1 and φ_2 are defined through

$$\tan \varphi_1 = (\Delta \beta_0 / \alpha) \tanh \alpha z \quad (6.51)$$

and

$$C_{12}^* = |C_{12}| e^{-j\varphi_2} \quad (6.52)$$

Although the total power $|a_1(z)|^2 - |a_2(z)|^2$ is kept constant, the transmission power carried by $a_1(z)$ alone increases with increasing z as shown in Fig. 6.4. If the coupling is removed at some point z , and $a_1(z)$ is fed into a load, the output power becomes $\{1 + (|C_{12}|/\alpha)^2 \sinh^2 \alpha z\}$ times the input power introduced at $z=0$. This type of power amplification by wave interaction can be realized in traveling wave tubes, which we shall discuss in Chapter 8.

6.4 Periodic Structures

In this section, we shall study spatially periodic structures, such as the one illustrated in Fig. 6.5. We can consider each period of the structure as a cavity and the whole structure as a chain of identical cavities. The electromagnetic field in each cavity interacts with those in the adjacent cavities in a complicated manner; however, a systematic study becomes possible if use is made of an eigenfunction approach. Although the input and output waveguides are aligned in Fig. 6.5, this is not essential in the present discussion. Here, all that is required is periodicity; the cavity between reference

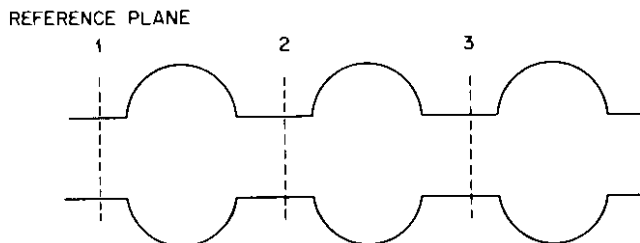


Fig. 6.5. A periodic structure.

planes 1 and 2 is identical with the cavity between 2 and 3, and so forth.

The transverse electric and magnetic fields at reference planes 1 and 2 can be expanded in terms of the waveguide eigenfunctions:

$$\begin{aligned} \mathbf{E}_{\parallel}(1) &= \sum \mathbf{E}_{in} V_n(1), & \mathbf{E}_{\parallel}(2) &= \sum \mathbf{E}_{in} V_n(2) \\ \mathbf{H}_{\parallel}(1) &= \sum \mathbf{k} \times \mathbf{E}_{in} I_n(1), & \mathbf{H}_{\parallel}(2) &= \sum \mathbf{k} \times \mathbf{E}_{in} I_n(2) \end{aligned} \quad (6.53)$$

where $V_n(i)$, $I_n(i)$ ($i = 1, 2$) are the expansion coefficients at reference plane i . Instead of taking an infinite number of terms with respect to n in the above summations, let us terminate at the N th term, then the field expressions become approximations. Since these approximations can be made as good as we desire by choosing sufficiently large N , a sacrifice in accuracy need not be considered in the following discussion. The transverse electric and magnetic fields are now uniquely determined by $\mathbf{v}_1$, $\mathbf{v}_2$, $\mathbf{i}_1$, and $\mathbf{i}_2$ which are defined by

$$\mathbf{v}_1 = \begin{bmatrix} V_1(1) \\ V_2(1) \\ \vdots \\ V_n(1) \end{bmatrix}, \quad \mathbf{v}_2 = \begin{bmatrix} V_1(2) \\ V_2(2) \\ \vdots \\ V_n(2) \end{bmatrix}, \quad \mathbf{i}_1 = \begin{bmatrix} I_1(1) \\ I_2(1) \\ \vdots \\ I_n(1) \end{bmatrix}, \quad \mathbf{i}_2 = \begin{bmatrix} -I_1(2) \\ -I_2(2) \\ \vdots \\ -I_n(2) \end{bmatrix} \quad (6.54)$$

The sign of $\mathbf{i}_2$ is reversed so that it represents currents flowing into the next cavity. Suppose $\mathbf{v}_1$ and $\mathbf{v}_2$ are somehow given, then the transverse electric fields at the input and output reference planes are fixed, and from them the electromagnetic field in the cavity can be calculated, in principle, as we did in Chapter 4. The transverse magnetic fields at the reference planes, and hence $\mathbf{i}_1$ and $\mathbf{i}_2$ can therefore generally be determined. Assuming the linearity of the field relation, this fact can be expressed in the form

$$\mathbf{i}_1 = \mathbf{Y}_{11} \mathbf{v}_1 + \mathbf{Y}_{12} \mathbf{v}_2, \quad -\mathbf{i}_2 = \mathbf{Y}_{21} \mathbf{v}_1 + \mathbf{Y}_{22} \mathbf{v}_2 \quad (6.55)$$

These equations indicate that if $\mathbf{v}_1$ and $\mathbf{v}_2$ are given, $\mathbf{i}_1$ and $\mathbf{i}_2$ can be determined. However, the expressions in (6.55) are not convenient for deriving an eigenvalue problem for the study of periodic structures since quantities at two different reference planes appear on both sides. We would like to place quantities related to one reference plane on one side and those related to the other reference plane on the other side. The cavity would then be represented by a matrix which transforms a vector representing quantities at one plane to another vector representing the same quantities at the other plane. Hence, n identical cavities in cascade could be represented simply by the n th power of the matrix; this is accomplished by defining a transfer-

matrix through the equation

$$\mathbf{w}_2 = \mathbf{T}\mathbf{w}_1 \quad (6.56)$$

where

$$\mathbf{w}_1 = \begin{bmatrix} \mathbf{v}_1 \\ \mathbf{i}_1 \end{bmatrix}, \quad \mathbf{w}_2 = \begin{bmatrix} \mathbf{v}_2 \\ \mathbf{i}_2 \end{bmatrix}, \quad \mathbf{T} = \begin{bmatrix} \mathbf{T}_{11} & \mathbf{T}_{12} \\ \mathbf{T}_{21} & \mathbf{T}_{22} \end{bmatrix} \quad (6.57)$$

By substitution, we see that

$$\begin{aligned} \mathbf{T}_{11} &= -\mathbf{Y}_{12}^{-1}\mathbf{Y}_{11}, & \mathbf{T}_{12} &= \mathbf{Y}_{12}^{-1} \\ \mathbf{T}_{21} &= -\mathbf{Y}_{21} + \mathbf{Y}_{22}\mathbf{Y}_{12}^{-1}\mathbf{Y}_{11}, & \mathbf{T}_{22} &= -\mathbf{Y}_{22}\mathbf{Y}_{12}^{-1} \end{aligned}$$

where an assumption is made that $\mathbf{Y}_{12}$ is nonsingular.

$\mathbf{T}$ is a square matrix of order $2N$, and hence it is expected that there are $2N$ independent eigenvectors; any vectors $\mathbf{w}$ can be expressed by their linear combination. Unless some restriction is imposed on $\mathbf{T}$, no fruitful discussion of the eigenvalue problem is possible; therefore, let us consider reciprocal and lossless conditions either of which can be used in the present discussion in the same way as (6.7) was used previously.

The reciprocal condition in terms of $\mathbf{T}$ is given by

$$\mathbf{T}_t \mathbf{R} \mathbf{T} = \mathbf{R} \quad (6.58)$$

where

$$\mathbf{R} = \begin{bmatrix} \mathbf{0} & -\mathbf{I} \\ \mathbf{I} & \mathbf{0} \end{bmatrix} \quad (6.59)$$

The proof is as follows. Rewriting (5.70) in terms of $\mathbf{v}_1$, $\mathbf{v}_2$, $\mathbf{i}_1$, and $\mathbf{i}_2$, we have

$$\mathbf{v}_{1t}^{(1)} \mathbf{i}_1^{(2)} - \mathbf{v}_{2t}^{(1)} \mathbf{i}_2^{(2)} - \mathbf{v}_{1t}^{(2)} \mathbf{i}_1^{(1)} + \mathbf{v}_{2t}^{(2)} \mathbf{i}_2^{(1)} = 0$$

Using (6.56) and (6.57), this can be expressed in terms of $\mathbf{v}_1$ and $\mathbf{i}_1$ only,

$$\begin{aligned} & \mathbf{v}_{1t}^{(1)} \mathbf{i}_1^{(2)} - (\mathbf{v}_{1t}^{(1)} \mathbf{T}_{11t} + \mathbf{i}_{1t}^{(1)} \mathbf{T}_{12t}) (\mathbf{T}_{21} \mathbf{v}_1^{(2)} + \mathbf{T}_{22} \mathbf{i}_1^{(2)}) \\ & - \mathbf{v}_{1t}^{(2)} \mathbf{i}_1^{(1)} + (\mathbf{v}_{1t}^{(2)} \mathbf{T}_{11t} + \mathbf{i}_{1t}^{(2)} \mathbf{T}_{12t}) (\mathbf{T}_{21} \mathbf{v}_1^{(1)} + \mathbf{T}_{22} \mathbf{i}_1^{(1)}) = 0 \end{aligned}$$

which is equivalent to

$$\begin{aligned} & \mathbf{v}_{1t}^{(1)} (\mathbf{I} - \mathbf{T}_{11t} \mathbf{T}_{22} + \mathbf{T}_{21t} \mathbf{T}_{12}) \mathbf{i}_1^{(2)} \\ & - \mathbf{v}_{1t}^{(2)} (\mathbf{I} - \mathbf{T}_{11t} \mathbf{T}_{22} + \mathbf{T}_{21t} \mathbf{T}_{12}) \mathbf{i}_1^{(1)} \\ & - \mathbf{v}_{1t}^{(1)} (\mathbf{T}_{11t} \mathbf{T}_{21} - \mathbf{T}_{21t} \mathbf{T}_{11}) \mathbf{v}_1^{(2)} \\ & - \mathbf{i}_{1t}^{(1)} (\mathbf{T}_{12t} \mathbf{T}_{22} - \mathbf{T}_{22t} \mathbf{T}_{12}) \mathbf{i}_1^{(2)} = 0 \end{aligned}$$

Noting that $\mathbf{v}_1^{(1)}$, $\mathbf{i}_1^{(1)}$, $\mathbf{v}_1^{(2)}$, and $\mathbf{i}_1^{(2)}$ are arbitrary, we obtain from this

relation

$$\begin{aligned}\mathbf{I} - \mathbf{T}_{11r}\mathbf{T}_{22} + \mathbf{T}_{21r}\mathbf{T}_{12} &= 0 \\ \mathbf{T}_{11r}\mathbf{T}_{21} - \mathbf{T}_{21r}\mathbf{T}_{11} &= 0 \\ \mathbf{T}_{12r}\mathbf{T}_{22} - \mathbf{T}_{22r}\mathbf{T}_{12} &= 0\end{aligned}$$

This completes the proof, since (6.58) is equivalent to these three equations.

The lossless condition in terms of $\mathbf{T}$ is given by

$$\mathbf{T}^+ \mathbf{L} \mathbf{T} = \mathbf{L} \quad (6.60)$$

where

$$\mathbf{L} = \begin{bmatrix} 0 & \mathbf{I} \\ \mathbf{I} & 0 \end{bmatrix} \quad (6.61)$$

The proof is as follows. If the circuit is lossless, the net power flowing into it must be equal to zero, i.e.,

$$\frac{1}{2} \{ (\mathbf{i}_1^+ \mathbf{v}_1 + \mathbf{v}_1^+ \mathbf{i}_1) - (\mathbf{i}_2^+ \mathbf{v}_2 + \mathbf{v}_2^+ \mathbf{i}_2) \} = 0$$

In terms of $\mathbf{w}_1$ and $\mathbf{w}_2$, this can be rewritten in the form

$$\mathbf{w}_1^+ \mathbf{L} \mathbf{w}_1 - \mathbf{w}_2^+ \mathbf{L} \mathbf{w}_2 = 0$$

which is equivalent to

$$\mathbf{w}_1^+ (\mathbf{L} - \mathbf{T}^+ \mathbf{L} \mathbf{T}) \mathbf{w}_1 = 0$$

where use is made of (6.56). Noting that this relation holds for arbitrary $\mathbf{w}_1$, we see that (6.60) is therefore proved.

We are now in a position to discuss the eigenvalue problem of $\mathbf{T}$, following the discussion given in Section 6.1. Let us first use (6.58).

The eigenvalue problem of $\mathbf{T}$ is given by

$$(\mathbf{T} - \gamma \mathbf{I}) \mathbf{x} = 0 \quad (6.62)$$

The eigenvalues of $\mathbf{T}$ are obtained by solving the algebraic equation for γ ,

$$\det(\mathbf{T} - \gamma \mathbf{I}) = 0 \quad (6.63)$$

If all the roots of this equation are distinct, Theorem I, in Section 6.1, and its proof holds for the eigenvectors of $\mathbf{T}$ without modification.

Corresponding to Theorem II, we find that if γ_k is an eigenvalue, γ_k^{-1} is also a eigenvalue. This can be proved as follows. First, we note that $\mathbf{T}$ is nonsingular and $\mathbf{T}^{-1}$ exists since $\det \mathbf{T} \neq 0$ as we can see by taking the determinant of (6.58) and using $\det \mathbf{R} \neq 0$. Applying $\mathbf{T}^{-1}$ from the right, (6.58) becomes

$$\mathbf{T}_l \mathbf{R} = \mathbf{R} \mathbf{T}^{-1} \quad (6.64)$$

Using this equation, we obtain the following relation from (6.63)

$$\begin{aligned}\det(\mathbf{T}_t - \gamma_k \mathbf{I}) \det \mathbf{R} &= \det(\mathbf{T}_t \mathbf{R} - \gamma_k \mathbf{R}) \\ &= \det(\mathbf{R} \mathbf{T}^{-1} - \gamma_k \mathbf{R}) = \det \mathbf{R} \det(\mathbf{T}^{-1} - \gamma_k \mathbf{I}) = 0\end{aligned}$$

Since $\det \mathbf{R} \neq 0$, we have

$$\det(\mathbf{T}^{-1} - \gamma_k \mathbf{I}) = 0$$

or equivalently,

$$\det(\mathbf{T} - \gamma_k^{-1} \mathbf{I}) = 0$$

which shows that γ_k^{-1} is also an eigenvalue of $\mathbf{T}$.

Corresponding to Theorem III, we find that if γ_k and γ_l are the eigenvalues corresponding to eigenvectors $\mathbf{x}_k$ and $\mathbf{x}_l$, respectively, and if $\gamma_k \neq \gamma_l^{-1}$, then

$$\mathbf{x}_{kt} \mathbf{R} \mathbf{x}_l = 0 \quad (6.65)$$

This can be shown as follows. Eigenvectors $\mathbf{x}_k$ and $\mathbf{x}_l$ satisfy

$$\mathbf{T} \mathbf{x}_k = \gamma_k \mathbf{x}_k, \quad \mathbf{T} \mathbf{x}_l = \gamma_l \mathbf{x}_l$$

which are equivalent to

$$\mathbf{x}_{kt} \mathbf{T}_t = \gamma_k \mathbf{x}_{kt}, \quad \mathbf{T}^{-1} \mathbf{x}_l = \gamma_l^{-1} \mathbf{x}_l$$

respectively. Multiplying the first equation by $\mathbf{R} \mathbf{x}_l$ from the right and the second equation by $\mathbf{x}_{kt} \mathbf{R}$ from the left, we obtain

$$\mathbf{x}_{kt} \mathbf{T}_t \mathbf{R} \mathbf{x}_l = \gamma_k \mathbf{x}_{kt} \mathbf{R} \mathbf{x}_l, \quad \mathbf{x}_{kt} \mathbf{R} \mathbf{T}^{-1} \mathbf{x}_l = \gamma_l^{-1} \mathbf{x}_{kt} \mathbf{R} \mathbf{x}_l$$

The left-hand sides of these two equations are equal from (6.64); therefore we have

$$(\gamma_k - \gamma_l^{-1}) \mathbf{x}_{kt} \mathbf{R} \mathbf{x}_l = 0$$

By hypothesis $\gamma_k \neq \gamma_l^{-1}$ and (6.65) is therefore deduced.

Finally, corresponding to Theorem IV, we find that if $\gamma_k = \gamma_l^{-1}$, then

$$\mathbf{x}_{kt} \mathbf{R} \mathbf{x}_l \neq 0 \quad (6.66)$$

From the first theorem, $(\mathbf{R} \mathbf{x}_l)^+$ can be expressed as a linear combination of eigenvectors.

$$(\mathbf{R} \mathbf{x}_l)^+ = \sum \alpha_m \mathbf{x}_{mt}$$

where the α_m 's are the expansion coefficients, and the eigenvectors are all transposed. Multiplying by $\mathbf{R} \mathbf{x}_l$ from the right, we have

$$(\mathbf{R} \mathbf{x}_l)^+ (\mathbf{R} \mathbf{x}_l) = \sum \alpha_m \mathbf{x}_{mt} \mathbf{R} \mathbf{x}_l$$

All the terms on the right-hand side disappear except the k th one, according to the previous theorem (6.65). The k th term cannot be equal to zero since the left-hand side is nonzero; this is impossible unless (6.66) holds, and the proof is therefore complete.

Let us define $\tilde{\mathbf{x}}_k$ by

$$\tilde{\mathbf{x}}_k = c\mathbf{x}_{kt} \quad (6.67)$$

when γ_k is either 1 or -1 , but when $\gamma_k^2 \neq 1$, then

$$\tilde{\mathbf{x}}_k = c\mathbf{x}_{lt} \quad (6.68)$$

where $\mathbf{x}_l$ is the eigenvector corresponding to $\gamma_l = \gamma_k^{-1}$ whose existence is guaranteed by the second theorem. By selecting c properly, let

$$\tilde{\mathbf{x}}_k \mathbf{R} \mathbf{x}_k = 1 \quad (6.69)$$

If $l \neq k$, we have from (6.65)

$$\tilde{\mathbf{x}}_l \mathbf{R} \mathbf{x}_k = 0 \quad (6.70)$$

The above discussion is based on the assumption that the $2N$ roots of (6.63) are all distinct. When (6.63) has multiple roots, we can still obtain $2N$ independent eigenvectors using a discussion similar to that presented in Section 6.1. Therefore, an arbitrary vector $\mathbf{x}$ in the $2N$ dimensional space can be expressed in the form

$$\mathbf{x} = \sum \mathbf{x}_k (\tilde{\mathbf{x}}_k \mathbf{R} \mathbf{x}) \quad (6.71)$$

where the summation is over k from 1 to $2N$. Multiplying by $\mathbf{T}$ from the left, we have

$$\mathbf{T} \mathbf{x} = \sum \gamma_k \mathbf{x}_k (\tilde{\mathbf{x}}_k \mathbf{R} \mathbf{x}) \quad (6.72)$$

The spectral representation of $\mathbf{T}$ is given by

$$\mathbf{T} \cdot = \sum \gamma_k \mathbf{x}_k \tilde{\mathbf{x}}_k \mathbf{R} \cdot \quad (6.73)$$

and that of $\mathbf{T}^n$ by

$$\mathbf{T}^n \cdot = \sum \gamma_k^n \mathbf{x}_k \tilde{\mathbf{x}}_k \mathbf{R} \cdot \quad (6.74)$$

The transverse electric and magnetic fields corresponding to an eigenvector of $\mathbf{T}$ are multiplied by the corresponding eigenvalue each time the reference plane is shifted one period in Fig. 6.5. Therefore, the electromagnetic field inside the n th period must be the n th power of the eigenvalue times the electromagnetic field inside the first period. Each independent electromagnetic field pattern with this property is called a mode of the periodic structure. It follows from the above discussion that any electromagnetic field in a periodic structure with reciprocity can be expressed as a linear

combination of the modes to any desired accuracy, assuming the existence of the transfer matrix.

Let us next use (6.60) in place of (6.58) to reach a similar conclusion. Assuming that the $2N$ roots of (6.63) are distinct, Theorem I holds without modification. Corresponding to Theorem II, we find that if γ_k is an eigenvalue, $(\gamma_k^*)^{-1}$ is also an eigenvalue of $\mathbf{T}$. To prove this, we note that $\mathbf{T}$ is nonsingular from (6.60), and hence the lossless condition can be written in the form

$$\mathbf{T}^+ \mathbf{L} = \mathbf{L} \mathbf{T}^{-1}$$

From this and (6.63), we obtain the following relation.

$$\begin{aligned} \det(\mathbf{T}^+ - \gamma_k^* \mathbf{I}) \det \mathbf{L} &= \det(\mathbf{T}^+ \mathbf{L} - \gamma_k^* \mathbf{L}) \\ &= \det(\mathbf{L} \mathbf{T}^{-1} - \gamma_k^* \mathbf{L}) = \det \mathbf{L} \det(\mathbf{T}^{-1} - \gamma_k^* \mathbf{I}) = 0 \end{aligned}$$

Since $\det \mathbf{L} \neq 0$, this shows that γ_k^* is an eigenvalue of $\mathbf{T}^{-1}$, or equivalently, $(\gamma_k^*)^{-1}$ is an eigenvalue of $\mathbf{T}$. Theorems III and IV become

$$\mathbf{x}_l^+ \mathbf{L} \mathbf{x}_k = 0 \quad (\gamma_l^{-1} \neq \gamma_k^*)$$

$$\mathbf{x}_l^+ \mathbf{L} \mathbf{x}_k \neq 0 \quad (\gamma_l^{-1} = \gamma_k^*)$$

The proofs are similar to those for the previous case, and we shall not repeat them. Once these relations are obtained, we can define $\tilde{\mathbf{x}}_k$ by

$$\tilde{\mathbf{x}}_k = c \mathbf{x}_k^+ \quad (|\gamma_k| = 1) \quad (6.75)$$

$$\tilde{\mathbf{x}}_k = c \mathbf{x}_l^+ \quad (|\gamma_k| \neq 1) \quad (6.76)$$

where $\mathbf{x}_l$ is the eigenvector corresponding to $\gamma_l = (\gamma_k^*)^{-1}$, whose existence is guaranteed by the second theorem. Let us choose c properly so that

$$\tilde{\mathbf{x}}_k \mathbf{L} \mathbf{x}_k = 1 \quad (6.77)$$

Then, using a discussion similar to the one given in Section 6.1, we have

$$\mathbf{x} = \sum \mathbf{x}_k (\tilde{\mathbf{x}}_k \mathbf{L} \mathbf{x}) \quad (6.78)$$

regardless of whether or not some of the eigenvalues are degenerate. The spectral representation of $\mathbf{T}$ is given by

$$\mathbf{T} \cdot = \sum \gamma_k \mathbf{x}_k \tilde{\mathbf{x}}_k \mathbf{L} \cdot \quad (6.79)$$

Modes in a lossless periodic structure can be defined in the same way as for the reciprocal case. We, therefore, conclude that any electromagnetic field in a lossless periodic structure can be expressed as a linear combination

of the modes to any desired accuracy, assuming the existence of the transfer matrix.

Each mode can be excited separately by generating a proper field pattern at a reference plane. Suppose that only one mode k is excited, then the surface integral $\int \mathbf{E} \times \mathbf{H}^* \cdot (-\mathbf{n}) dS$ over reference planes 1 and 2 ($\mathbf{n}$ points outwards) becomes

$$(\mathbf{i}_1^+ \mathbf{v}_1 - \mathbf{i}_2^+ \mathbf{v}_2) = \mathbf{i}_1^+ \mathbf{v}_1 (1 - \gamma_k \gamma_k^*)$$

When $|\gamma_k|^2 = 1$, no attenuation takes place between one reference plane and the next one, and the mode is said to be in the passband. In this case, the above surface integral vanishes and (2.84) shows that the average electric and magnetic energies stored in one spatial period are equal, assuming nonzero ω . When $|\gamma_k|^2 \neq 1$, since $\gamma_k \neq (\gamma_k^*)^{-1}$, we have

$$\mathbf{x}_k^+ \mathbf{L} \mathbf{x}_k = 0$$

or equivalently,

$$[\mathbf{v}_1^+ \mathbf{i}_1^+] \begin{bmatrix} 0 & \mathbf{I} \\ \mathbf{I} & 0 \end{bmatrix} \begin{bmatrix} \mathbf{v}_1 \\ \mathbf{i}_1 \end{bmatrix} = \mathbf{i}_1^+ \mathbf{v}_1 + \mathbf{v}_1^+ \mathbf{i}_1 = 0$$

Therefore, no real power is transmitted, and the mode is said to be in the stopband.

Let us next apply (5.96) to one spatial period in Fig. 6.5. To do so, we assume both the reciprocal and lossless conditions. For the k th mode, the left-hand side becomes

$$\begin{aligned} & \mathbf{i}_1^+ \frac{\partial \mathbf{v}_1}{\partial \omega} + \mathbf{v}_1^+ \frac{\partial \mathbf{i}_1}{\partial \omega} - \mathbf{i}_2^+ \frac{\partial \mathbf{v}_2}{\partial \omega} - \mathbf{v}_2^+ \frac{\partial \mathbf{i}_2}{\partial \omega} \\ &= \left(\mathbf{i}_1^+ \frac{\partial \mathbf{v}_1}{\partial \omega} + \mathbf{v}_1^+ \frac{\partial \mathbf{i}_1}{\partial \omega} \right) (1 - \gamma_k^* \gamma_k) + (\mathbf{i}_1^+ \mathbf{v}_1 + \mathbf{v}_1^+ \mathbf{i}_1) \left(-\gamma_k^* \frac{\partial \gamma_k}{\partial \omega} \right) \end{aligned}$$

In the passband, the first term on the right-hand side vanishes, since $|\gamma_k|^2 = 1$, and (5.96) reduces to

$$-\gamma_k^* (\partial \gamma_k / \partial \omega) = \frac{1}{2} j P^{-1} \int (\mu \mathbf{H}^* \cdot \mathbf{H} + \epsilon \mathbf{E}^* \cdot \mathbf{E}) dv \quad (6.80)$$

where P is the transmission power given by

$$P = \operatorname{Re} \{ \mathbf{i}_1^+ \mathbf{v}_1 \}$$

Let L be the length of one period and β be defined by

$$\gamma_k = \exp(-j\beta L)$$

Then (6.80) becomes

$$(\partial\beta/\partial\omega) = \frac{1}{2}P^{-1}L^{-1} \int (\mu\mathbf{H}^* \cdot \mathbf{H} + \epsilon\mathbf{E}^* \cdot \mathbf{E}) dv \quad (6.81)$$

Comparing this with (3.23) in the discussion of waveguides, we see that the left-hand side corresponds to the reciprocal of the group velocity. From this observation, we may say that the group velocity of a mode in the passband is equal to the transmission-power times the length of one period divided by the electromagnetic energy stored therein.

In a lossless, reciprocal, periodic structure, if γ_k is an eigenvalue, γ_k^{-1} and $(\gamma_k^*)^{-1}$ must also be eigenvalues. If γ_n is not real and $|\gamma_k| \neq 1$, then at least four different eigenvalues $\gamma_k, \gamma_k^*, \gamma_k^{-1}, (\gamma_k^*)^{-1}$ have to exist. When we take into account only one $\mathbf{E}_{in}$ in our calculation, as we shall do in the next section, $\mathbf{T}$ becomes a two-by-two matrix, and hence only two eigenvalues can exist. In this case, either the eigenvalues have to be real or their magnitude is unity, otherwise four different eigenvalues are required.

6.5 ω - β Diagram

Let us consider a waveguide with identical windows periodically spaced as shown in Fig. 6.6, and suppose that only one waveguide mode is necessary to approximate the fields at the reference planes located midway between any two adjacent windows. Let b_0 represent the normalized susceptance of the window with respect to the waveguide admittance. Then an equivalent circuit representing one period of the structure becomes a transmission line with electrical length $\beta_0 L$ shunted by an admittance $jb_0 Y_0$ at the midpoint as shown in Fig. 6.7. The term β_0 is the phase constant and Y_0 is the charac-

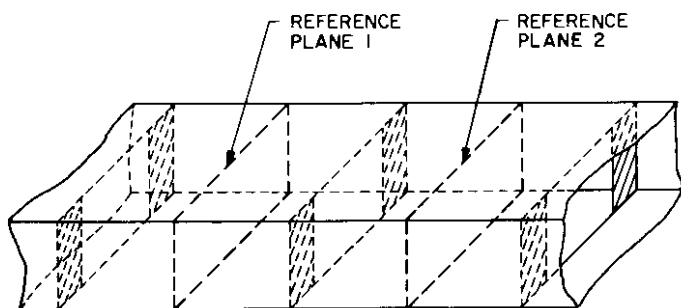


Fig. 6.6. Waveguide with periodically spaced windows.

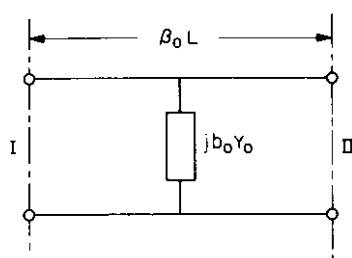


Fig. 6.7. Equivalent circuit for one section of the periodically loaded waveguide.

teristic admittance of the waveguide mode under consideration, and L is the length of one period.

The voltage and current at port I can be expressed in terms of the incident and reflected waves A and B (see Section 3.5) as follows:

$$V_1 = A + B, \quad I_1 = Y_0(A - B) \quad (6.82)$$

The voltage and current at the left-hand side of $j b_0 Y_0$ are given by

$$V = A e^{-j\theta} + B e^{j\theta}, \quad I = Y_0(A e^{-j\theta} - B e^{j\theta}) \quad (6.83)$$

where $\theta = \beta_0 L/2$. The voltage and current at the right-hand side of $j b_0 Y_0$ become V and $I - j b_0 Y_0 V$, respectively, since current $j b_0 Y_0 V$ is diverted into the shunt admittance. In terms of the waves A' and B' on the right-hand side of the window, V and $I - j b_0 Y_0 V$ must be expressible in the forms

$$V = A' + B', \quad I - j b_0 Y_0 V = Y_0(A' - B') \quad (6.84)$$

and the voltage and current at port II are given by

$$V_2 = A' e^{-j\theta} + B' e^{j\theta}, \quad I_2 = Y_0(A' e^{-j\theta} - B' e^{j\theta}) \quad (6.85)$$

Note that the direction of I_2 is chosen to be outward so as to conform with the sign of i_2 in Section 6.4.

Eliminating V and I from (6.83) and (6.84), A' and B' can be expressed in terms of A and B . Substituting this result into (6.85) and then using (6.82), we can write V_2 and I_2 in terms of V_1 and I_1 . The result is given by

$$\begin{bmatrix} V_2 \\ I_2 \end{bmatrix} = \mathbf{T} \begin{bmatrix} V_1 \\ I_1 \end{bmatrix} \quad (6.86)$$

where $\mathbf{T}$ is a square matrix of order 2 with the following components:

$$\begin{aligned} T_{11} &= T_{22} = \cos \beta_0 L - \frac{1}{2} b_0 \sin \beta_0 L \\ T_{12} &= j Y_0^{-1} \left(\frac{1}{2} b_0 - \sin \beta_0 L - \frac{1}{2} b_0 \cos \beta_0 L \right) \\ T_{21} &= -j Y_0 \left(\frac{1}{2} b_0 + \sin \beta_0 L + \frac{1}{2} b_0 \cos \beta_0 L \right) \end{aligned} \quad (6.87)$$

The two eigenvalues of $\mathbf{T}$, γ_1 and γ_2 , are obtained by solving the quadratic equation $\det(\mathbf{T} - \gamma \mathbf{I}) = 0$.

$$\gamma_{1,2} = \cos \beta_0 L - \frac{1}{2} b_0 \sin \beta_0 L \mp \{ (\frac{1}{2} b_0)^2 - (\sin \beta_0 L + \frac{1}{2} b_0 \cos \beta_0 L)^2 \}^{1/2} \quad (6.88)$$

When

$$(\frac{1}{2} b_0)^2 < (\sin \beta_0 L + \frac{1}{2} b_0 \cos \beta_0 L)^2 \quad (6.89)$$

the eigenvalues become complex and their magnitudes are found to be unity, indicating that the corresponding modes are in the passband. On the other hand, if

$$(\frac{1}{2} b_0)^2 > (\sin \beta_0 L + \frac{1}{2} b_0 \cos \beta_0 L)^2 \quad (6.90)$$

both of the eigenvalues are real and their magnitudes are not equal to unity, which means the corresponding modes are in the stopband.

If the opening of a waveguide window is small, b_0 is large and the passband condition (6.89) is satisfied only when $|\cos \beta_0 L|$ is close to unity; or equivalently, when the length of one period of the structure becomes approximately an integer multiple of a half-wavelength, thereby forming a resonator. When the length deviates from such a value, the sections between two adjacent windows cease to resonate, and no transmission of power becomes possible thus resulting in a stopband. Similarly, if b_0 is small, the stopband condition (6.90) is satisfied only when the length of one period becomes approximately an integer multiple of a half-wavelength. In this case, although the reflections from individual windows are small, the reflections from successive windows arrive back in phase and the overall result is no transmission of power. When the length of one period deviates appreciably from an integer multiple of a half-wavelength, the small reflections do not add in phase and introduce only a minor modification in the wave propagation.

Let us now concentrate on the passband. Since one of the eigenvalues is the complex conjugate of the other and their magnitudes are unity, γ_1 and γ_2 can be expressed as

$$\gamma_1 = e^{-j\beta L}, \quad \gamma_2 = e^{j\beta L} \quad (6.91)$$

where the phase shift from one reference plane to the next is indicated by

βL . Substituting (6.91) into (6.88), we obtain

$$\cos \beta L = \cos \beta_0 L - \frac{1}{2} b_0 \sin \beta_0 L \quad (6.92)$$

from which we learn how βL changes with ω .

From (3.127), the relation between β_0 and ω is given by

$$\beta_0^2 = \omega^2 \epsilon \mu - k_0^2$$

where k_0^2 is the eigenvalue of the waveguide mode. Therefore, ω as a function of $\beta_0 L$ should appear as the broken line in Fig. 6.8(a). If b_0 is a slowly varying function of ω , the right-hand side of (6.92) as a function of

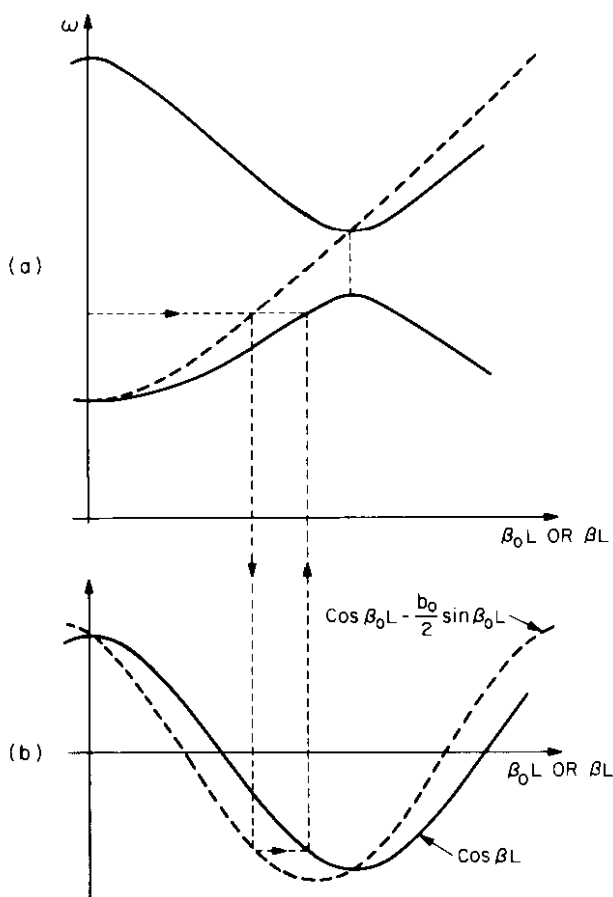


Fig. 6.8. Construction of the curve ω versus βL .

$\beta_0 L$ may, look like the broken line in Fig. 6.8(b). The solid line indicates $\cos \beta L$ as a function of βL . For a given ω , $\beta_0 L$ is determined from Fig. 6.8(a), and hence the corresponding value of βL can be found from Fig. 6.8(b) following the dotted lines with arrowheads. Thus, ω versus βL should appear as the solid lines in Fig. 6.8(a). Note that the broken line ω versus $\beta_0 L$ and one of the solid lines ω versus βL intersect at $\beta L = \beta_0 L = n\pi$ (n : integer). Also note that the tangent to the solid lines becomes horizontal at $\beta L = n\pi$, indicating that the group velocity becomes zero at the corresponding ω 's. Those ranges of ω for which there is no corresponding βL in Fig. 6.8(a) represent the stopbands.

A graph which shows ω versus β in a similar way is called the ω - β diagram for the periodic structure. From the above discussion, we see that the general appearance of the ω - β diagram for a periodic structure should look like Fig. 6.9. Each of the curves extends from $\beta = -\infty$ to $+\infty$ as a periodic function of β with a period $2\pi/L$. Since βL and $\beta L + 2n\pi$ give the same eigenvalues in (6.91), one period of the curves, say from $\beta = -\pi/L$ to π/L , contains all the information about the eigenvalues, and the remaining periods are redundant. The particular range of β mentioned above is called the first Brillouin zone. As far as the phase shift of the electromagnetic field at the reference planes is concerned, only the first Brillouin zone is of interest. The situation is slightly different, however, if we are interested in the more general problem of the fields throughout the waveguide. Let us consider the electric field, first concentrating on the mode corresponding to γ_1 . Each time we move to the next section of the periodic structure, the phase of the electric field changes by βL since the phases of the electromagnetic fields at the boundary reference planes shift the same amount. Thus, we have

$$\mathbf{E}(x, y, z + nL) = e^{-j\beta nL} \mathbf{E}(x, y, z)$$

where the z -axis is taken in the longitudinal direction of the waveguide. The above expression is equivalent to

$$e^{j\beta z} \mathbf{E}(x, y, z) = e^{j\beta(z+nL)} \mathbf{E}(x, y, z + nL)$$

which indicates that $e^{j\beta z} \mathbf{E}(x, y, z)$ is a periodic function of z with period L . Using a Fourier series expansion, it is expressed in the form

$$e^{j\beta z} \mathbf{E}(x, y, z) = \sum_{n=-\infty}^{\infty} \mathbf{A}_n(x, y) \exp(j2n\pi z/L)$$

or equivalently, we have

$$\mathbf{E}(x, y, z) = \sum_{n=-\infty}^{\infty} \mathbf{A}_n(x, y) \exp[-j\{\beta - (2n\pi/L)\}z]$$

In other words, the electric field consists of many components, each with the same group velocity but different phase velocities; the velocity of the n th component is given by

$$v_p = \omega \{\beta - (2n\pi/L)\}^{-1}$$

The n th component is sometimes called the n th space harmonic. The phase velocity appears as the slope of the dotted line drawn from the origin to the corresponding point on the ω - β diagram as shown in Fig. 6.9. Depending

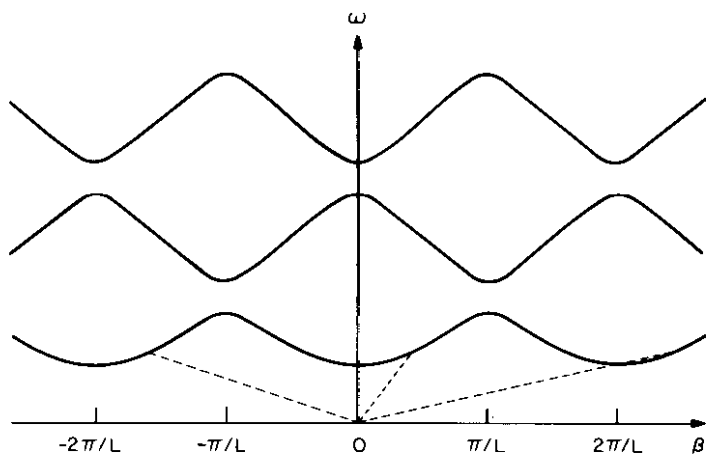


Fig. 6.9. An example of an ω - β diagram.

on the value of n , the phase velocity becomes very small, and for some components it may even become negative; these components appear to travel in the direction opposite to the power transmission. The magnetic field corresponding to γ_1 and the electric and magnetic fields corresponding to γ_2 can be discussed in a similar manner.

In many electron beam devices, such as the traveling wave tube which we shall discuss in Chapter 8, the necessary interaction between the circuit and the electron beam requires a slow electric field component traveling with a velocity approximately equal to that of the electrons. Since various periodic structures are used to obtain such a slow component, periodic structures

are often called slow-wave structures. In the traveling wave tube, the structure must be designed so as to make the z -component of $\mathbf{A}_n(x, y)$ at the electron beam as large as possible and the other field components as small as possible in order to minimize unnecessary losses. It should be noted that $\mathbf{A}_n(x, y)$ has the proper phase velocity.

In the above analysis, an assumption was made that only one waveguide mode was necessary to approximate the fields at the reference planes. As ω increases, however, more and more waveguide modes start to propagate and sooner or later the assumption becomes invalid. Several propagating modes may interact with each other in a complicated manner, and the ω - β diagram fails to resemble the one derived in this section. To avoid such a complicated situation, the operating frequency is generally kept low and/or the structure is made to satisfy certain symmetry conditions under which the excitation of higher order propagating modes is prohibited.

PROBLEMS

- 6.1 Following the argument in Section 6.3., discuss the exchange of transmission power between the two coupled modes under the condition that $a_1(z) = A_0$ and $a_2(z) = -A_0$ at $z = 0$.
- 6.2 In the above problem change the condition at $z = 0$ to $a_1(z) = a_2(z) = A_0$.
- 6.3 Give an example of a periodic structure without $\mathbf{Y}_{ij}$ ($i, j = 1, 2$). Also give an example with $\mathbf{Y}_{ij}$ but with a singular $\mathbf{Y}_{12}$.
- 6.4 When only one propagating mode is excited in a periodic structure, the average electric and magnetic energies stored in one period are equal to each other as shown in Section 6.4. Investigate the cases in which two modes are excited, and give an example in which the electric and magnetic energies stored in one period are not equal to each other.
- 6.5 Prove that when several modes are excited in a lossless periodic structure, the total transmission power is equal to the sum of the transmission powers due to the individual modes.
- 6.6 Suppose that each period of a periodic structure has a plane of symmetry at its midpoint and obtain the condition that must be satisfied by $\mathbf{T}$. Also investigate what kinds of relations exist among the eigenvalues.
- 6.7 Construct the ω - β diagram for an ordinary transmission line (TEM mode) having identical shunt capacitances periodically spaced.

CHAPTER 7

LINEAR AMPLIFIERS

In many practical applications, a signal level may be too low to perform a desired operation without introducing excessive noise or distortion. In such cases the signal is often amplified to a more useful level. To indicate the amount of amplification introduced by the amplifier, we define the transducer gain as the ratio of the actual signal power absorbed in the load to the available signal power from the generator. An amplifier is supposed to have a transducer gain larger than unity. When the transducer gain is large over a wide frequency range, the amplifier is described as having a broadband and high gain.

If excessive noise or distortion is introduced during the amplification process, the amplifier may become useless, even though it provides high gain. Consequently, the noise performance and distortion are considered to be important factors in determining the quality of amplifiers. In this chapter, we shall concentrate mostly on the amplifier noise performance, assuming that the input-output relation is linear, i.e., neglecting any possible distortion. This does not mean that distortion is unimportant; on the contrary, in many cases it becomes a vital factor. Unfortunately, no unified theory of amplifier distortion is available in a form suitable for presentation here.

The conventional theory of amplifier noise performance almost exclusively express results in terms of the available gain, which is defined as the ratio of available signal power at the amplifier output port to the available signal power from the generator. This quantity is convenient because the available gain of a cascade amplifier is equal to the product of the available gains of the component amplifiers. In addition, the discussion of noise performance

of the cascade amplifier can be made particularly simple since the use of available gain conceals the real complexity which exists. On the other hand, when the gain of practical amplifiers is discussed, it is generally in terms of their transducer gain. If the available gain is used, the available signal power at the output port is of primary interest rather than the actual signal power, and hence the user is left with the burden of providing a proper load impedance to utilize it; this often turns out to be the most difficult part of the whole operation. Both the generator and load impedances are therefore specified beforehand, and the amplifier is designed to deliver the amplified signal to the specified load impedance, i.e., to provide a proper transducer gain. Usually, the generator and load impedances are equal to the characteristic impedances of the transmission lines at the amplifier input and output ports, respectively. Under these circumstances, if we develop the theory of amplifier noise performance with emphasis on the available gain, many confusing or misleading conclusions are reached. For this reason, we shall follow an alternative course and develop a theory based on the transducer gain.

We first derive an equivalent circuit for linear noisy multiport networks. The noise performance of amplifiers is discussed next, considering an amplifier as a two-port, linear, noisy network with a transducer gain larger than unity. The unconditional stability conditions and the unilateral gain are also presented. The last three sections are devoted to the discussion of tunnel-diode and parametric amplifiers that illustrate the kind of noise performance to be expected from these practical devices.

7.1 Canonical Forms

It was shown in Section 5.3 that the exchangeable power of a generator is invariant to lossless nonsingular transformations. A different proof of this theorem will now be given, using the impedance matrix of a lossless two-port network since it will set a pattern for deriving invariants of n -port noisy networks.

In the circuit shown in Fig. 7.1, the exchangeable power of the generator connected to port I of the two-port network is given by

$$P_0 = |E|^2 / (4 \operatorname{Re} Z_i) \quad (7.1)$$

and the relation among V_1 , V_2 , I_1 , and I_2 is given by

$$\begin{bmatrix} V_1 \\ V_2 \end{bmatrix} = \begin{bmatrix} Z_{11} & Z_{12} \\ Z_{21} & Z_{22} \end{bmatrix} \begin{bmatrix} I_1 \\ I_2 \end{bmatrix} \quad (7.2)$$

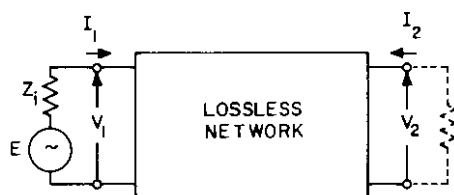


Fig. 7.1. Lossless transformation of a one-port generator.

Since V_1 is equal to $E - Z_i I_1$, we have from (7.2)

$$V_1 = Z_{11} I_1 + Z_{12} I_2 = E - Z_i I_1 \quad (7.3)$$

Assuming that $Z_{11} + Z_i$ is not equal to zero, I_1 is obtained in terms of E and I_2 ,

$$I_1 = (Z_{11} + Z_i)^{-1} (E - Z_{12} I_2)$$

Substituting this into (7.2), we have

$$V_2 = Z_{21} I_1 + Z_{22} I_2 = Z_{21} (Z_i + Z_{11})^{-1} E + \{Z_{22} - Z_{12} Z_{21} (Z_i + Z_{11})^{-1}\} I_2 \quad (7.4)$$

Comparing this with (7.3), and remembering that the positive direction of I_2 is inward, the open-circuit voltage and the output impedance looking back at port II are given by

$$Z_{21} (Z_i + Z_{11})^{-1} E \quad \text{and} \quad \{Z_{22} - Z_{12} Z_{21} (Z_i + Z_{11})^{-1}\}$$

respectively. The exchangeable power at port II therefore becomes

$$P'_0 = |E|^2 |Z_{21}|^2 |Z_i + Z_{11}|^{-2} / 4 \operatorname{Re} \{Z_{22} - Z_{12} Z_{21} (Z_i + Z_{11})^{-1}\} \quad (7.5)$$

the denominator of which can be rewritten in the form

$$\begin{aligned} & 4 \operatorname{Re} \{Z_{22} - Z_{12} Z_{21} (Z_i + Z_{11})^{-1}\} \\ &= 2 \left\{ Z_{22} + Z_{22}^* - \frac{Z_{12} Z_{21}}{Z_i + Z_{11}} - \frac{Z_{12}^* Z_{21}^*}{(Z_i + Z_{11})^*} \right\} \end{aligned} \quad (7.6)$$

When the two-port network is lossless, the impedance matrix Z must satisfy the lossless condition

$$Z + Z^+ = 0$$

or equivalently,

$$Z_{11} = -Z_{11}^*, \quad Z_{12} = -Z_{21}^*, \quad Z_{22} = -Z_{22}^*$$

With these relations, the right-hand side of (7.6) reduces to

$$2 \frac{|Z_{21}|^2 (Z_i + Z_{11} + Z_i^* + Z_{11}^*)}{|Z_i + Z_{11}|^2} = |Z_{21}|^2 |Z_i + Z_{11}|^{-2} 4 \operatorname{Re} Z_i$$

Replacing the denominator of (7.5) by this expression, we have

$$P_0' = |E|^2 / (4 \operatorname{Re} Z_i) = P_0 \quad (7.7)$$

which shows the exchangeable power is invariant to the lossless transformation. It is worth noting that we assumed nonzero $Z_{11} + Z_i$ and nonzero Z_{21} in the above calculation. The assumption of nonzero Z_{21} was used to avoid making (7.5) indeterminate.

Since the exchangeable power has been defined only for one-port generators, a question arises whether or not there are some quantities for an n -port generator which are invariant to lossless transformations and hence correspond to the exchangeable power of a one-port generator. To answer this question, let us study a lossless transformation of a linear n -port generator following the pattern set by the above discussion.

Some voltages may appear at each port of the n -port generator when all the ports are open circuited. Let $E_1, E_2, \dots, E_n$ be the voltages where the subscripts refer to the port numbers. We define $\mathbf{e}$ by

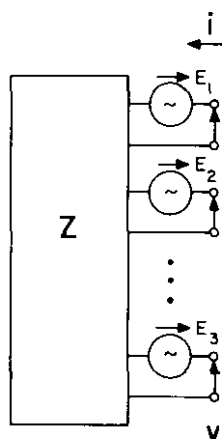
$$\mathbf{e} = \begin{bmatrix} E_1 \\ E_2 \\ \vdots \\ E_n \end{bmatrix} \quad (7.8)$$

Referring again to the discussion concerning an equivalent circuit for one-port generators in Section 1.4, we see that the relation between the voltages and current at the ports is given by

$$\mathbf{v} = \mathbf{Z}\mathbf{i} + \mathbf{e} \quad (7.9)$$

where $\mathbf{v}$ is the voltage vector, $\mathbf{i}$ the current vector, and $\mathbf{Z}$ the internal impedance matrix of the n -port generator. An equivalent circuit for (7.9) is shown in Fig. 7.2.

Since we used a two-port lossless network to transform a one-port generator, we should consider using a $2n$ -port lossless network to transform an n -port generator. Taking the voltage and current directions for the $2n$ -

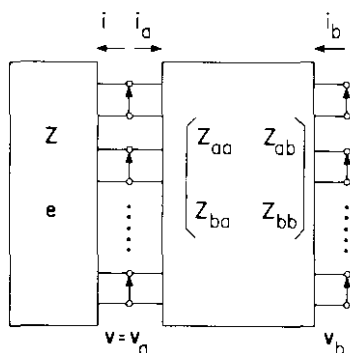
Fig. 7.2. Equivalent circuit of an n -port generator.

port network as shown in Fig. 7.3, we have

$$\begin{bmatrix} \mathbf{v}_a \\ \mathbf{v}_b \end{bmatrix} = \begin{bmatrix} \mathbf{Z}_{aa} & \mathbf{Z}_{ab} \\ \mathbf{Z}_{ba} & \mathbf{Z}_{bb} \end{bmatrix} \begin{bmatrix} \mathbf{i}_a \\ \mathbf{i}_b \end{bmatrix} \quad (7.10)$$

where $\mathbf{v}_a$ and $\mathbf{v}_b$ are the voltage vectors, while $\mathbf{i}_a$ and $\mathbf{i}_b$ are the current vectors at reference planes a and b , respectively. The terms $\mathbf{Z}_{aa}$, $\mathbf{Z}_{ab}$, $\mathbf{Z}_{ba}$, and $\mathbf{Z}_{bb}$ are all square matrices of order n , and together they form the impedance matrix of the $2n$ -port network. Since $\mathbf{v} = \mathbf{v}_a$, we have from (7.9) and (7.10)

$$\mathbf{Z}\mathbf{i} + \mathbf{e} = \mathbf{Z}_{aa}\mathbf{i}_a + \mathbf{Z}_{ab}\mathbf{i}_b$$

Fig. 7.3. Lossless transformation of an n -port generator.

Keeping in mind that $\mathbf{i}$ and $\mathbf{i}_a$ are equal, except for their directions being opposite, we obtain from the above equation

$$\mathbf{i}_a = (\mathbf{Z} + \mathbf{Z}_{aa})^{-1} (\mathbf{e} - \mathbf{Z}_{ab}\mathbf{i}_b)$$

where an assumption is made that $(\mathbf{Z} + \mathbf{Z}_{aa})$ is nonsingular. Substituting this into (7.10), $\mathbf{v}_b$ can be expressed in terms of $\mathbf{e}$ and $\mathbf{i}_b$;

$$\begin{aligned} \mathbf{v}_b &= \mathbf{Z}_{ba}\mathbf{i}_a + \mathbf{Z}_{bb}\mathbf{i}_b \\ &= \{\mathbf{Z}_{bb} - \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1}\mathbf{Z}_{ab}\}\mathbf{i}_b + \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1}\mathbf{e} \end{aligned} \quad (7.11)$$

Let $\mathbf{e}'$ be the open-circuited voltage vector and $\mathbf{Z}'$ the internal impedance matrix observed at reference plane b , then we have from (7.11)

$$\mathbf{e}' = \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1}\mathbf{e} \quad (7.12)$$

$$\mathbf{Z}' = \mathbf{Z}_{bb} - \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1}\mathbf{Z}_{ab} \quad (7.13)$$

The exchangeable power of a one-port generator is given by the square of the absolute value of the open-circuited voltage divided by four times the real part of the internal impedance. Corresponding to the real part of the internal impedance, we should take one half of $\mathbf{Z} + \mathbf{Z}^+$, the sum of the impedance matrix and its adjoint. Corresponding to the square of the magnitude of the open-circuited voltage, $|E|^2 = EE^*$, we may consider either $\mathbf{ee}^+$ or $\mathbf{e}^+\mathbf{e}$; however, since $\mathbf{e}^+\mathbf{e}$ becomes a matrix of order 1×1 , or equivalently, a scalar quantity, it does not match the above matrix form of the impedance. Thus, $\mathbf{ee}^+$, a square matrix of order n , seems to be a natural choice. In following sections, we shall study the case in which the E_i 's represent noise voltages, and hence the ensemble average $\langle E_i E_j^* \rangle$ is more meaningful than the particular value of $E_i E_j^*$ itself. For this reason, let us choose $\langle \mathbf{ee}^+ \rangle$, whose ij component is $\langle E_i E_j^* \rangle$, instead of $\mathbf{ee}^+$ itself. Corresponding to the fact that $|E|^2$ is either positive or zero, $\langle \mathbf{ee}^+ \rangle$ is either positive-definite or positive-semidefinite. This can be seen from

$$\langle \mathbf{x}^+ \mathbf{e} (\mathbf{x}^+ \mathbf{e})^+ \rangle = \mathbf{x}^+ \langle \mathbf{ee}^+ \rangle \mathbf{x} \geq 0 \quad (7.14)$$

where $\mathbf{x}$ is an arbitrary fixed vector and the angle brackets $\langle \rangle$ indicate the ensemble average of the quantity within.

We are now in a position to investigate with the help of (7.12) and (7.13) how $\langle \mathbf{ee}^+ \rangle$ and $\mathbf{Z} + \mathbf{Z}^+$ are transformed by the lossless $2n$ -port network. From the lossless condition, the sum of the impedance matrix of the $2n$ -

port network and its adjoint must be zero;

$$\begin{bmatrix} \mathbf{Z}_{aa} & \mathbf{Z}_{ab} \\ \mathbf{Z}_{ba} & \mathbf{Z}_{bb} \end{bmatrix} + \begin{bmatrix} \mathbf{Z}_{aa}^+ & \mathbf{Z}_{ba}^+ \\ \mathbf{Z}_{ab}^+ & \mathbf{Z}_{bb}^+ \end{bmatrix} = 0$$

or equivalently,

$$\mathbf{Z}_{aa} = -\mathbf{Z}_{aa}^+, \quad \mathbf{Z}_{ab} = -\mathbf{Z}_{ba}^+, \quad \mathbf{Z}_{bb} = -\mathbf{Z}_{bb}^+ \quad (7.15)$$

From this, we see that $\mathbf{Z}_{aa}$ and $\mathbf{Z}_{bb}$ are restricted by the condition that their adjoint matrices must be equal to themselves except for the opposite signs. On the other hand, there is no real restriction imposed on $\mathbf{Z}_{ba}$. Once $\mathbf{Z}_{ba}$ is given, $\mathbf{Z}_{ab}$ is fixed by the middle condition of (7.15); it therefore follows that either $\mathbf{Z}_{ba}$ or $\mathbf{Z}_{ab}$ can be chosen arbitrarily.

For convenience, let us define $\mathbf{H}^+$ as

$$\mathbf{H}^+ = \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1} \quad (7.16)$$

Since $\mathbf{Z}_{ba}$ is arbitrary, $\mathbf{H}^+$ can also be chosen arbitrarily by selecting a proper lossless network for the transformation. Using $\mathbf{H}^+$ thus defined, $\langle \mathbf{e}'\mathbf{e}'^+ \rangle$ can be written in the form

$$\langle \mathbf{e}'\mathbf{e}'^+ \rangle = \langle \mathbf{H}^+ \mathbf{e}(\mathbf{H}^+ \mathbf{e})^+ \rangle = \mathbf{H}^+ \langle \mathbf{e}\mathbf{e}^+ \rangle \mathbf{H} \quad (7.17)$$

where (7.12) is used. This equation indicates how $\langle \mathbf{e}\mathbf{e}^+ \rangle$ is transformed by the $2n$ -port network. For the transformation of $\mathbf{Z} + \mathbf{Z}^+$, let us consider $\mathbf{Z}' + \mathbf{Z}'^+$. From (7.13), we have

$$\mathbf{Z}' + \mathbf{Z}'^+ = \mathbf{Z}_{bb} - \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1} \mathbf{Z}_{ab} + \mathbf{Z}_{bb}^+ - \mathbf{Z}_{ab}^+(\mathbf{Z}^+ + \mathbf{Z}_{aa}^+)^{-1} \mathbf{Z}_{ba}^+$$

Using (7.15), this can be simplified as follows:

$$\begin{aligned} \mathbf{Z}' + \mathbf{Z}'^+ &= \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1} \mathbf{Z}_{ba}^+ + \mathbf{Z}_{ba}(\mathbf{Z}^+ + \mathbf{Z}_{aa}^+)^{-1} \mathbf{Z}_{ba}^+ \\ &= \mathbf{Z}_{ba} \{ (\mathbf{Z} + \mathbf{Z}_{aa})^{-1} + (\mathbf{Z}^+ + \mathbf{Z}_{aa}^+)^{-1} \} \mathbf{Z}_{ba}^+ \\ &= \mathbf{Z}_{ba}(\mathbf{Z} + \mathbf{Z}_{aa})^{-1} \{ (\mathbf{Z}^+ + \mathbf{Z}_{aa}^+) + (\mathbf{Z} + \mathbf{Z}_{aa}) \} (\mathbf{Z}^+ + \mathbf{Z}_{aa}^+)^{-1} \mathbf{Z}_{ba}^+ \\ &= \mathbf{H}^+ (\mathbf{Z} + \mathbf{Z}^+) \mathbf{H} \end{aligned} \quad (7.18)$$

A comparison of (7.17) and (7.18) shows that $\langle \mathbf{e}\mathbf{e}^+ \rangle$ and $\mathbf{Z} + \mathbf{Z}^+$ undergo the same transformation by the lossless $2n$ -port network. Since $(\mathbf{Z} + \mathbf{Z}^+)^+ = \mathbf{Z} + \mathbf{Z}^+$ and $(\langle \mathbf{e}\mathbf{e}^+ \rangle)^+ = \langle \mathbf{e}\mathbf{e}^+ \rangle$, it follows that $\langle \mathbf{e}\mathbf{e}^+ \rangle$ and $\mathbf{Z} + \mathbf{Z}^+$ are both self-adjoint. Furthermore, $\langle \mathbf{e}\mathbf{e}^+ \rangle$ is either positive-definite or positive-semidefinite from (7.14). Let us assume for the moment that $\langle \mathbf{e}\mathbf{e}^+ \rangle$ is positive-definite. Then there are two self-adjoint matrices, and one of them is positive-definite. Consequently, it is possible to select a nonsingular $\mathbf{H}$ which diagonalizes both $\langle \mathbf{e}'\mathbf{e}'^+ \rangle$ and $\mathbf{Z}' + \mathbf{Z}'^+$, simultaneously, as

shown by (5.60) and (5.61). The diagonal components of $\mathbf{Z}' + \mathbf{Z}'^+$ are given by the eigenvalues of $(\langle \mathbf{e}\mathbf{e}^+ \rangle)^{-1} (\mathbf{Z} + \mathbf{Z}^+)$ which are all real. Let $\lambda_1, \lambda_2, \dots, \lambda_n$ be the real eigenvalues, then we have

$$\langle \mathbf{e}'\mathbf{e}'^+ \rangle = \mathbf{I} \quad (7.19)$$

$$\mathbf{Z}' + \mathbf{Z}'^+ = \text{diag} [\lambda_1 \quad \lambda_2 \quad \dots \quad \lambda_n] \quad (7.20)$$

In order for $\mathbf{H}$ to exist, $\mathbf{Z} + \mathbf{Z}_{aa}$ has to be nonsingular, and for $\mathbf{H}$ to be nonsingular, $\mathbf{Z}_{ba}$ has to be nonsingular. Since these conditions do not contradict the lossless condition given by (7.15), it is always possible to get an appropriate lossless $2n$ -port network which gives (7.19) and (7.20).

Let us investigate the physical meaning of the result obtained above. It follows from (7.20) that the new n -port generator has an internal impedance matrix of the form

$$\mathbf{Z}' = \frac{1}{2} \text{diag} [\lambda_1 \quad \lambda_2 \quad \dots \quad \lambda_n] + \mathbf{Z}_r \quad (7.21)$$

where $\mathbf{Z}_r$ is an n -port impedance matrix satisfying the lossless condition

$$\mathbf{Z}_r + \mathbf{Z}_r^+ = 0 \quad (7.22)$$

In order to obtain (7.19) and (7.20), we have selected an appropriate $\mathbf{H}$; however, since $\mathbf{H}$ contains only $\mathbf{Z}_{ba}$ and $\mathbf{Z}_{aa}$, $\mathbf{Z}_{bb}$ of the $2n$ -port network can still be selected arbitrarily provided that $\mathbf{Z}_{bb}^+ + \mathbf{Z}_{bb} = 0$. Since the subtraction of $\mathbf{Z}_r$ from $\mathbf{Z}_{bb}$ does not change this condition, we can always select a proper $\mathbf{Z}_{bb}$ which makes $\mathbf{Z}_r$ in (7.21) disappear. Remembering that $\lambda_1, \lambda_2, \dots, \lambda_n$ are all real, this means that by selecting a proper lossless $2n$ -port network the transformed circuit is equivalent to n resistors, each connected in series with a voltage source expressed by a component of $\mathbf{e}'$. Since $\langle \mathbf{e}'\mathbf{e}'^+ \rangle = \mathbf{I}$, each voltage source is uncorrelated with the others and has a unit magnitude as illustrated in box C on the left-hand side of Fig. 7.4. This is called the canonical form of the original n -port generator, and the operation performed by the lossless $2n$ -port network is called the canonical transformation.

The lossless $2n$ -port network transforms $\mathbf{v}$ and $\mathbf{i}$ of the n -port generator into $\mathbf{v}_b$ and $\mathbf{i}_b$. Since $\mathbf{v} = \mathbf{v}_a$ and $\mathbf{i} = -\mathbf{i}_a$, we have

$$\mathbf{v} = -\mathbf{Z}_{aa}\mathbf{i} + \mathbf{Z}_{ab}\mathbf{i}_b \quad (7.23)$$

$$\mathbf{v}_b = -\mathbf{Z}_{ba}\mathbf{i} + \mathbf{Z}_{bb}\mathbf{i}_b \quad (7.24)$$

From (7.23), $\mathbf{i}_b$ is given by

$$\mathbf{i}_b = \mathbf{Z}_{ab}^{-1}\mathbf{v} + \mathbf{Z}_{ab}^{-1}\mathbf{Z}_{aa}\mathbf{i}$$

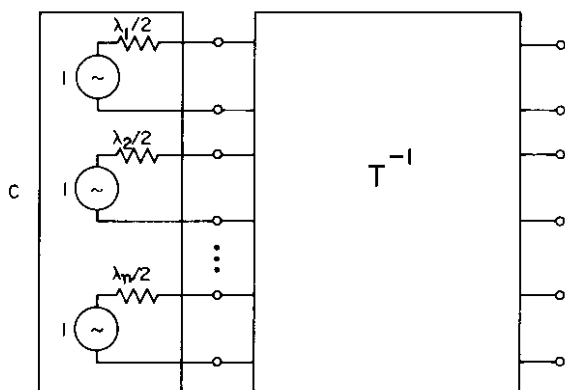


Fig. 7.4. Equivalent circuit of an n -port generator including the canonical form C and $2n$ -port transforming network T^{-1} where T is the canonical transformation.

where the existence of Z_{ab}^{-1} is guaranteed since H , and hence Z_{ba} , is chosen to be nonsingular and $Z_{ab} = -Z_{ba}^+$ from (7.15). Substituting the above equation into (7.24), we obtain

$$v_b = Z_{bb}Z_{ab}^{-1}v + (Z_{bb}Z_{ab}^{-1}Z_{aa} - Z_{ba})i$$

The relation between the voltages and currents at the input and output ports of the $2n$ -port network is therefore given by

$$\begin{bmatrix} v_b \\ -i_b \end{bmatrix} = \begin{bmatrix} Z_{bb}Z_{ab}^{-1} & Z_{ba} - Z_{bb}Z_{ab}^{-1}Z_{aa} \\ -Z_{ab}^{-1} & Z_{ab}^{-1}Z_{aa} \end{bmatrix} \begin{bmatrix} v \\ -i \end{bmatrix} = T \begin{bmatrix} v \\ -i \end{bmatrix} \quad (7.25)$$

Thus defined, T is the transfer matrix representing the lossless $2n$ -port network which performs the canonical transformation; T satisfies the lossless condition (6.60), and if T exists, it is nonsingular from (6.60). In general, the transformation performed by a circuit with nonsingular transfer matrix is said to be nonsingular; therefore the canonical transformation is nonsingular.

Multiplying (6.60) by $(T^+)^{-1}$ from the left and by T^{-1} from the right, we obtain

$$(T^{-1})^+ L T^{-1} = L \quad (7.26)$$

A comparison of (7.26) with (6.60) shows that T^{-1} represents another lossless $2n$ -port network. Suppose that the original n -port generator is first

transformed by $\mathbf{T}$ into its canonical form and is next transformed by $\mathbf{T}^{-1}$, then the resultant network has exactly the same $\mathbf{v}$ and $\mathbf{i}$ as those of the original n -port generator. We, therefore, conclude that an arbitrary n -port generator is equivalent to its canonical form followed by a lossless $2n$ -port network whose transfer matrix is given by $\mathbf{T}^{-1}$.

The eigenvalues, $\lambda_1, \lambda_2, \dots, \lambda_n$, of $(\langle \mathbf{e}\mathbf{e}^+ \rangle)^{-1} (\mathbf{Z} + \mathbf{Z}^+)$ are invariant to a nonsingular, lossless, but otherwise arbitrary, transformation $\mathbf{T}_1$. This follows because the cascade connection of two lossless $2n$ -port networks with transfer matrices $\mathbf{T}_1^{-1}$ and $\mathbf{T}$ gives the same canonical form characterized by $\lambda_1, \lambda_2, \dots, \lambda_n$. Furthermore, since the canonical form is solely determined by the eigenvalues, there are no invariants other than $\lambda_1, \lambda_2, \dots, \lambda_n$ and their functions when arbitrary nonsingular lossless transformations are considered.

For one-port generators, the exchangeable power is invariant to lossless nonsingular transformations. This is understandable because the exchangeable power is given by $1/(2\lambda_1)$, where λ_1 is the eigenvalue of $(\langle \mathbf{e}\mathbf{e}^+ \rangle)^{-1} \times (\mathbf{Z} + \mathbf{Z}^+)$ which is a matrix of order 1×1 in this case. The exchangeable power and its functions are the only invariants to lossless nonsingular transformations.

Finally, let us investigate the case in which $\langle \mathbf{e}\mathbf{e}^+ \rangle$ is positive-semidefinite. In this case $(\langle \mathbf{e}\mathbf{e}^+ \rangle)^{-1}$ does not exist; however, by adding a small voltage vector $\Delta \mathbf{e}$ to the open-circuited voltage $\mathbf{e}$ at the terminals $\langle (\mathbf{e} + \Delta \mathbf{e}) \times (\mathbf{e} + \Delta \mathbf{e})^+ \rangle$ can be made positive-definite. For this perturbed n -port generator, the canonical form can be obtained as we have done above. In this canonical form, $\lambda_1, \lambda_2, \dots, \lambda_n$ are functions of the small voltages making up the vector $\Delta \mathbf{e}$. Since the perturbed generator approaches the original generator as $\Delta \mathbf{e}$ approaches zero, let us assume that the canonical form of the original n -port generator is obtained by taking the limit as $\Delta \mathbf{e} \rightarrow 0$.

During this limiting process, one or more of $\lambda_1, \lambda_2, \dots, \lambda_n$ may become infinite showing that the values of the corresponding resistor are infinite. There are two possible cases: (i) The resistor should have a finite value in series with a zero-voltage source, but since a unit voltage source is assumed, the resistance becomes infinite; (ii) The resistor is actually open-circuited and its value is infinite. It is difficult to distinguish these two cases in a general discussion.

Summarizing the above results, the following conclusion will be reached: An n -port generator is equivalent to n resistors each with an independent voltage source, followed by a lossless $2n$ -port transforming network as illustrated in Fig. 7.4. When $\langle \mathbf{e}\mathbf{e}^+ \rangle$ is positive-semidefinite, the number of the resistors may be less than n and/or some of the voltage sources disappear.

7.2 Noise Measure

An amplifier can be considered as a two-port network which has specific generator and load impedances assigned to it. In addition the network generally contains noise sources. From the conclusion reached in the previous section, such a two-port network is equivalent to the canonical form followed by a lossless four-port network. The canonical form may have two resistors, one resistor, or none at all. With a passive network, the output power can never be larger than the input power, hence at least one resistor in the canonical form of an amplifier must be negative. This enables us to concentrate on the following three cases: The canonical form has (i) two negative resistances, (ii) one negative resistance and one positive resistance, (iii) only one negative resistance.

Let us first study case (i) for which the equivalent amplifier circuit is illustrated in Fig. 7.5. Ports I and II of the lossless four-port network

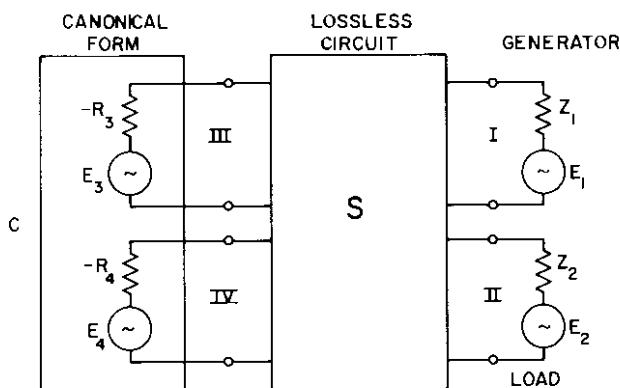


Fig. 7.5. Equivalent circuit of a linear noisy amplifier.

represent the input and output ports of the amplifier, respectively. Ports III and IV are connected to the negative resistances $Z_3 = -R_3$ and $Z_4 = -R_4$ in the canonical form. Let incident and reflected waves a and b be defined such that their i th components are given by

$$a_i = \frac{1}{2} |\operatorname{Re} Z_i|^{-1/2} (V_i + Z_i I_i), \quad b_i = \frac{1}{2} |\operatorname{Re} Z_i|^{-1/2} (V_i - Z_i^* I_i) \quad (7.27)$$

where the subscript refers to the port number. The scattering matrix S of

the lossless four-port network times $\mathbf{a}$ then becomes $\mathbf{b}$,

$$\mathbf{b} = \mathbf{S}\mathbf{a} \quad (7.28)$$

From (5.87), $\mathbf{S}$ satisfies the lossless condition

$$\mathbf{S}^+ \mathbf{P} \mathbf{S} = \mathbf{P} \quad (7.29)$$

where $\mathbf{P}$ is given by

$$\mathbf{P} = \text{diag}[1 \quad 1 \quad -1 \quad -1] \quad (7.30)$$

since negative resistances are connected to ports III and IV. Multiplying (7.29) by $(\mathbf{P}\mathbf{S})^{-1} = \mathbf{S}^{-1}\mathbf{P}^{-1}$ from the right and by $\mathbf{S}\mathbf{P}^{-1}$ from the left, we have

$$\mathbf{S}\mathbf{P}^{-1}\mathbf{S}^+ = \mathbf{P}^{-1} \quad (7.31)$$

Equations (7.29) and (7.31) are equivalent to each other, and either one can be used as the lossless condition for $\mathbf{S}$. The 22 component of (7.31) gives

$$|S_{21}|^2 + |S_{22}|^2 - |S_{23}|^2 - |S_{24}|^2 = 1 \quad (7.32)$$

which is the condition the S_{2i} 's must satisfy.

With this much preparation, we are in a position to discuss the noise performance of the amplifier in detail. The transducer gain of an amplifier is defined by

$$G = \frac{\text{actual signal power into load}}{\text{available signal power from generator}} \quad (7.33)$$

For an amplifier shown in Fig. 7.5, the numerator is given by the signal component of $|b_2|^2$. If $|a_1|^2$ represents the available signal power of the generator, $|b_2|^2 = |S_{21}|^2 |a_1|^2$. Thus, we have

$$G = |b_2|^2 / |a_1|^2 = |S_{21}|^2 \quad (7.34)$$

The exchangeable power is not used in the denominator of (7.33) because when the real part of the generator impedance is negative, the actual signal power flowing into a load impedance with a positive real part can be made as large as one wishes by inserting a lossless nonamplifying network.

The operating noise temperature T_{op} of an amplifier is defined by

$$T_{op} = \frac{\text{actual noise power into load}}{kBG} \quad (7.35)$$

where k is the Boltzmann constant, and the noise power is measured in a narrow bandwidth B with a generator having a noise temperature T_i .

From the definition of noise temperature, the available noise power of the generator is given by kT_iB . When $B = 1$ MHz and $T_i = 290^\circ\text{K}$, kT_iB is 4×10^{-15} W, or equivalently, -114 dBm. In general, the available noise power from an ordinary resistor at $T^\circ\text{K}$ is given by kTB as derived by a statistical consideration and confirmed by experiments, which we shall not discuss any further.

The noise originating in and reflected back to the load is just as troublesome as other noise components, and hence it is included in the operating noise temperature of (7.35). This can be seen from the definition of actual power given in Section 1.4. When the load impedance has a positive real part, the actual noise power flowing into the load is defined to be the net noise power transferred to the load plus the available noise power of the load, all in the bandwidth B . Thus, the operating noise temperature T_{op} is a function of the noise temperature T_L of the load, and T_L may be assumed to be 290°K if not stated otherwise.

For the amplifier shown in Fig. 7.5, the ratio of the operating noise temperature T_{op} to the generator noise temperature T_i is given by

$$\frac{T_{op}}{T_i} = \frac{|b_2|^2}{|a_1|^2 G} = \frac{|S_{21}|^2 |a_1|^2 + |S_{22}|^2 |a_2|^2 + |S_{23}|^2 |a_3|^2 + |S_{24}|^2 |a_4|^2}{|S_{21}|^2 |a_1|^2} \quad (7.36)$$

where all the power waves are comprised of noise components only and $|a_1|^2 = kT_iB$.

Using G and the noise temperature ratio T_{op}/T_i , we define the operating noise measure of an amplifier as

$$M = \frac{(T_{op}/T_i) - 1}{1 - (1/G)} \quad (7.37)$$

If the transducer gain is large, M is essentially equal to $(T_{op} - T_i)/T_i$ which is the ratio of the noise introduced in the amplification process to the amplified noise from the generator, $(kT_{op}BG - kT_iBG)/kT_iBG$. A negative value of M means that G is smaller than unity. When G is larger than unity, as it should be for an amplifier, M is always positive; the amplifier will obviously become less noisy as M becomes smaller. When G decreases and $(T_{op}/T_i) - 1$ does not become smaller in value, then M cannot remain unchanged. When the gain is small, the added noise must be correspondingly small for the amplifier to be considered equally noisy.

For a comparison of different amplifiers we use $M \times T_i$, rather than M itself, for the following reason: If T_i is increased for one amplifier, keeping

everything else the same, $|S_{21}|^2 |a_1|^2$ increases in (7.36) and the output of the amplifier becomes noisier. At the same time M decreases giving the misleading impression that the amplifier is now quieter. On the other hand, $M \times T_i$ is independent of T_i , indicating the noise introduced in the amplification process divided by $kB(G-1)$.

Those who are familiar with the IRE (Institute of Radio Engineers) definition of noise figure F will notice that for a high gain amplifier with $T_i = 290^\circ\text{K}$ and a negligible load noise contribution, M is essentially equal to the excess noise figure $(F-1)$ of the amplifier.

The noise figure F , at a specified input frequency, is defined as: the ratio of (1) the total noise power per unit bandwidth at a corresponding output frequency "available" at the output port when the noise temperature of the input termination is standard (290°K) to (2) that portion of (1) engendered at the input frequency by the input termination. The contribution from the noise originating in and reflected back to the load is not included in this definition. Realizing its importance, however, the IRE subsequently defined the operating noise temperature so as to include the noise contribution from the load.

In our case, M can be written in the form

$$M = \frac{|S_{22}|^2 |a_2|^2 + |S_{23}|^2 |a_3|^2 + |S_{24}|^2 |a_4|^2}{(|S_{21}|^2 - 1) |a_1|^2} \quad (7.38)$$

where (7.34) and (7.36) are used. Rewriting the denominator with the help of (7.32), (7.38) becomes

$$M = \frac{|S_{22}|^2 |a_2|^2 + |S_{23}|^2 |a_3|^2 + |S_{24}|^2 |a_4|^2}{(-|S_{22}|^2 + |S_{23}|^2 + |S_{24}|^2) |a_1|^2} \quad (7.39)$$

The expression inside the parentheses in the denominator cannot become less than -1 since $|S_{21}|^2 \geq 0$ in (7.32); $|S_{22}|^2$, $|S_{23}|^2$, and $|S_{24}|^2$ otherwise can take any positive values.

Without a loss of generality, assume that $|a_3|^2 \leq |a_4|^2$ and under this assumption, let us investigate the possible range of the value of M . When $|S_{22}|^2 > |S_{23}|^2 + |S_{24}|^2$, M is negative. In addition, the condition

$$|S_{23}|^2 = |S_{24}|^2 = 0$$

gives the minimum value for $|M|$, because when $|S_{23}|^2$ or $|S_{24}|^2$ departs from zero, the denominator on the right-hand side of (7.39) becomes smaller in magnitude and the numerator larger. Thus, M cannot come closer to

zero than $-|a_2|^2/|a_1|^2$ in the negative range. Since $|a_2|^2$ is given by $kT_L B$ and $|a_1|^2$ by $kT_i B$, the maximum value in the negative range is

$$M_m = -|a_2|^2/|a_1|^2 = -(T_L/T_i) \quad (7.40)$$

When $|S_{22}|^2 < |S_{23}|^2 + |S_{24}|^2$, M is positive. If $|S_{23}|^2$ and $|S_{24}|^2$ are kept constant, the condition $|S_{22}|^2 = 0$ minimizes M since the denominator in (7.39) decreases and the numerator increases if $|S_{22}|^2$ departs from zero. Therefore, in order to obtain the smallest value of M in the positive range, we have only to consider the case in which $|S_{22}|^2 = 0$. Since $|a_3|^2 \leq |a_4|^2$ by hypothesis, we have

$$\frac{|S_{23}|^2 |a_3|^2 + |S_{24}|^2 |a_4|^2}{(|S_{23}|^2 + |S_{24}|^2) |a_1|^2} \geq \frac{|S_{23}|^2 |a_3|^2 + |S_{24}|^2 |a_3|^2}{(|S_{23}|^2 + |S_{24}|^2) |a_1|^2} = \frac{|a_3|^2}{|a_1|^2}$$

It follows that when $|S_{22}|^2 = |S_{24}|^2 = 0$, M attains the smallest positive value M_{opt} which is called the optimum noise measure and is given by

$$M_{\text{opt}} = |a_3|^2/|a_1|^2 = -(P_{e, \min}/kT_i B) \quad (7.41)$$

where $P_{e, \min}$ indicates the exchangeable noise power with the smaller magnitude obtainable from the two negative resistances in the equivalent circuit. Later we shall discuss a case in which many negative resistances are involved.

The magnitude of M can increase indefinitely in the negative as well as in the positive range if the scattering coefficients are properly selected. The solid lines in Fig. 7.6, therefore, give the possible range of M . In order to obtain M_{opt} , the conditions $|S_{22}|^2 = 0$ and $|S_{24}|^2 = 0$ must be satisfied,

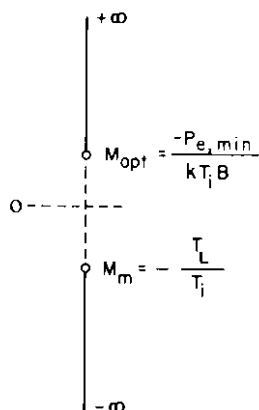


Fig. 7.6. The possible range of M .

unless $T_L = 0$ and $|a_3|^2 = |a_4|^2$, respectively. In other words, the output port must satisfy the matching condition and the negative resistance with the smaller absolute value in the canonical form must be effectively disconnected from the output port.

Given an amplifier, one method to achieve M_{opt} is shown in Fig. 7.7.

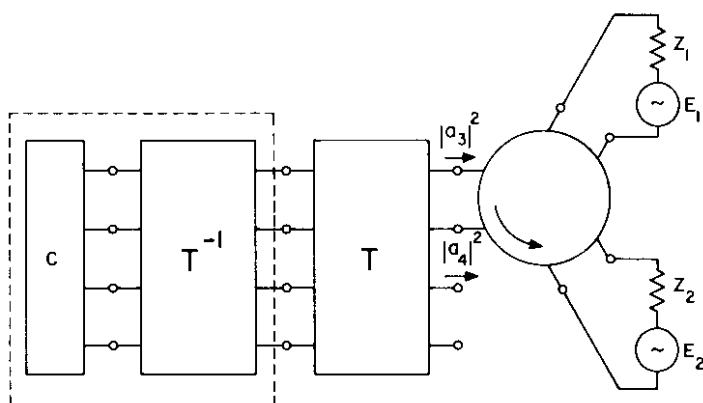


Fig. 7.7. Realization of M_{opt} .

Here the amplifier is represented by an equivalent circuit comprising the canonical form C followed by a lossless transformation T^{-1} . The lossless four-port network with the transfer matrix T connected to the amplifier restores the canonical form and the negative resistance with exchangeable noise power $P_{e, \min} = -|a_3|^2$ in the canonical form is connected to the generator and load through a three-port circulator. The other resistance is left open. Both conditions $|S_{24}|^2 = 0$ and $|S_{22}|^2 = 0$ are satisfied by this arrangement, and hence M_{opt} is realized. It may be worth mentioning that the input port is also matched in this realization of M_{opt} .

We have so far discussed case (i) in which two resistances in the canonical form are negative. The discussions for the other two cases are similar. When one resistance is negative and the other is positive in the canonical form, we assume that the negative one is connected to port III and the positive one to port IV of the lossless four-port network shown in Fig. 7.5. Then, corresponding to (7.39), we have

$$M = \frac{|S_{22}|^2 |a_2|^2 + |S_{23}|^2 |a_3|^2 + |S_{24}|^2 |a_4|^2}{(-|S_{22}|^2 + |S_{23}|^2 - |S_{24}|^2) |a_1|^2}$$

When there is only one negative resistance in the canonical form, we assume

that the negative resistance is connected to port III, and that port IV is eliminated from the beginning. Then, we have

$$M = \frac{|S_{22}|^2 |a_2|^2 + |S_{23}|^2 |a_3|^2}{(-|S_{22}|^2 + |S_{23}|^2) |a_1|^2}$$

For both cases, the optimum noise measure is given by

$$M_{\text{opt}} = |a_3|^2 / |a_1|^2 = -(P_{e, \text{min}} / kT_i B)$$

where $P_{e, \text{min}}$ is equal to the exchangeable noise power of the negative resistance in each case.

Let us now ask the following question: What is the best (i.e., smallest positive) noise measure obtainable when we are given several amplifiers

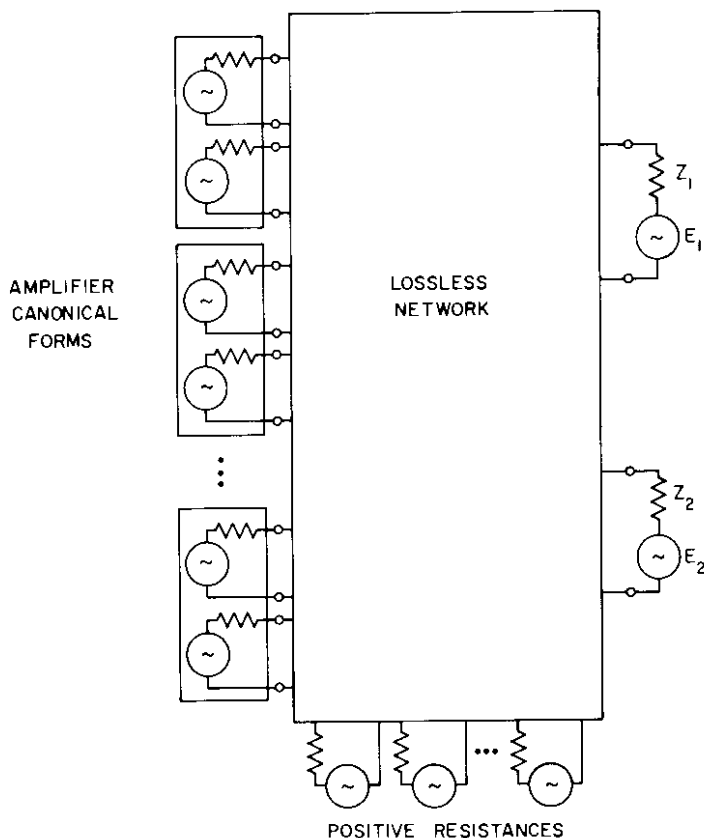


Fig. 7.8. Equivalent circuit of the composite amplifier.

and are allowed to use any passive networks in order to connect them? A passive network can be considered to be a lossless multiport network with some of the ports connected to positive resistances. The most general equivalent circuit of the amplifier we can construct is therefore given by Fig. 7.8 where the lossless parts of the component amplifiers are all included in the lossless network. Once this equivalent circuit is obtained, we can easily apply an argument similar to the one previously used, and the optimum noise measure is given by

$$M_{\text{opt}} = -(P_{e, \text{min}}/kT_i B)$$

where $P_{e, \text{min}}$ is the exchangeable noise power with the smallest magnitude among all the negative resistances in the equivalent circuit. This result shows that there is no possibility of improving the noise measure by combining amplifiers. The best value we can obtain from the combined amplifier is equal to the best value obtainable from the single amplifier having the smallest M_{opt} .

On the other hand, M_m is given by

$$M_m = -(T_{\text{min}}/T_i)$$

where T_{min} is the lowest noise temperature of all the positive resistances in the equivalent circuit. The range of possible M is from M_{opt} through infinity to M_m , which is similar to Fig. 7.6.

Suppose that two amplifiers are connected in cascade through a circulator as shown in Fig. 7.9. The third port of the circulator is terminated by the matched resistance R_0 having a noise temperature T_{L1} . The reference impedances of the circulator used to define the scattering matrix are assumed to be given by

$$Z_1 = Z_{L1}^*, \quad Z_2 = Z_{g2}^*, \quad Z_3 = R_0 \quad (7.42)$$

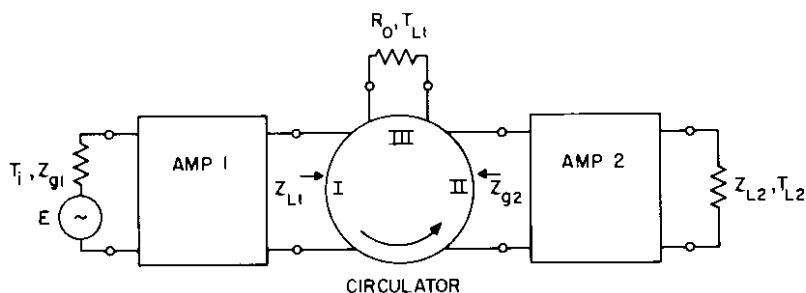


Fig. 7.9. Cascade connection of two amplifiers.

so that both amplifiers see the correct load and generator impedances assigned to them. Since the actual power flowing into port I of the circulator is transformed to the available power from port II, the overall gain G becomes the product of the individual amplifier gains:

$$\begin{aligned}
 G &= \frac{\text{actual signal power into load}}{\text{available signal power from generator}} \\
 &= \frac{\text{actual signal power into circulator}}{\text{available signal power from generator}} \\
 &\quad \times \frac{\text{actual signal power into load}}{\text{available signal power from circulator}} \\
 &= G_1 G_2
 \end{aligned} \tag{7.43}$$

where G_1 and G_2 represent the transducer gains of the first and second amplifiers, respectively.

The actual noise power N flowing into the load is given by the sum of the amplified noise and the noise added by the second amplifier. The added noise is equal to $M_2 k T_i B (G_2 - 1)$, where M_2 is the noise measure of the second amplifier with generator noise temperature T_i . The available noise power at the input of the second amplifier is equal to the actual noise power flowing into port I of the circulator in Fig. 7.9, namely, $M_1 k T_i B (G_1 - 1) + k T_i B G_1$, where M_1 is the noise measure of the first amplifier with generator noise temperature T_i . Thus, we have

$$N = k T_i B G_1 G_2 + M_1 k T_i B (G_1 - 1) G_2 + M_2 k T_i B (G_2 - 1)$$

from which the overall operating noise temperature T_{op} can be obtained as follows:

$$T_{op} = \frac{N}{k B G} = T_i \left(1 + M_1 - M_1 \frac{1}{G_1} + M_2 \frac{G_2 - 1}{G_1 G_2} \right) \tag{7.44}$$

where use is made of (7.43). Substituting (7.43) and (7.44) into (7.37), the overall noise measure becomes

$$\begin{aligned}
 M &= \left(M_1 - M_1 \frac{1}{G_1} + M_2 \frac{G_2 - 1}{G_1 G_2} \right) \left(1 - \frac{1}{G_1 G_2} \right)^{-1} \\
 &= M_1 + (M_2 - M_1) \frac{G_2 - 1}{G_1 G_2 - 1}
 \end{aligned} \tag{7.45}$$

Suppose that the two amplifiers are identical, then $M_1 = M_2$ and the second term on the right-hand side of (7.45) disappears which makes $M = M_1$ and

$G = G_1^2$. If we connect another amplifier to this cascaded amplifier through a second circulator, we obtain an amplifier with $M = M_1$ and $G = G_1^3$. Continuing this process, if we connect a large number of identical amplifiers in cascade, we obtain a high gain amplifier having the same noise measure as that of the individual amplifiers. From this observation, we can interpret the noise measure of an individual amplifier as the excess noise figure $(F_\infty - 1)$ of a high gain amplifier constructed by cascading a large number of such amplifiers. When the amplifier under consideration is matched at both ends and $Z_g = Z_L^*$, the circulators can be eliminated from the above discussion.

One might think that the circulator with a resistor R_0 at the third port can be replaced by a lossless network in Fig. 7.9. This is generally not true, however. First of all, the noise temperature of the input impedance of the second amplifier may not be the same as that of the load impedance of the first amplifier. Furthermore, if the first amplifier has an output impedance with a negative real part, there is no way of transforming this to have a positive real part by any lossless transformation, and hence the second amplifier is unable to see a proper generator impedance. The circulator effectively isolates the two amplifiers which would otherwise interfere with each other as described.

The discussion in this section can be summarized as follows:

- (i) The smaller the positive value of M is, the less noisy the amplifier.
- (ii) Negative noise measure means that the gain of the amplifier is less than one, i.e., attenuation occurs rather than amplification.
- (iii) M_{opt} can be realized by a lossless transformation.
- (iv) It is impossible to achieve positive M less than M_{opt} by any passive transformation.
- (v) The M_{opt} of a combined amplifier is equal to the best M_{opt} among the component amplifiers.

7.3 Unconditional Stability

An amplifier is said to be unconditionally stable if the real parts of its input and output impedances remain positive when the load and generator impedances are changed arbitrarily provided their real parts remain positive. Let us consider the condition for the unconditional stability. Let S_{11} , S_{12} , S_{21} , and S_{22} represent the scattering coefficients of the amplifier where subscripts 1 and 2 refer to the input and output ports, respectively. We assume

that the real parts of the reference impedances Z_1 and Z_2 are both positive and that $|S_{11}| < 1$ and $|S_{22}| < 1$; otherwise the real parts of the input and output impedances cannot be assumed to remain positive. Suppose that the load impedance Z_L is connected to the output port, then a_2 is found to be

$$a_2 = \{(Z_L - Z_2)/(Z_L + Z_2^*)\} b_2 = r_2 b_2$$

in a similar manner to the derivation of (5.132) where r_2 is the reflection coefficient of Z_L with respect to Z_2^* . Since we keep the real part of Z_L positive, $|r_2| < 1$. The equation $b = Sa$ becomes

$$b_1 = S_{11}a_1 + S_{12}r_2b_2, \quad b_2 = S_{21}a_1 + S_{22}r_2b_2$$

from which we obtain

$$b_1 = \{S_{11} + S_{12}r_2S_{21}(1 - r_2S_{22})^{-1}\} a_1 \quad (7.46)$$

It follows from this that the condition for the real part of the input impedance to remain positive is given by

$$|S_{11} + S_{12}r_2S_{21}(1 - r_2S_{22})^{-1}| < 1$$

which is equivalent to

$$|-AS_{22}^{-1} + B(1 - r_2S_{22})^{-1}| < 1 \quad (7.47)$$

where

$$A = S_{12}S_{21} - S_{11}S_{22}, \quad B = S_{12}S_{21}S_{22}^{-1} \quad (7.48)$$

The problem has now been reduced to finding the condition for (7.47) to hold regardless of the value of r_2 when $|r_2| < 1$.

To find the above condition, let us next investigate the bilinear transformation

$$z = -AS_{22}^{-1} + B(1 - r_2S_{22})^{-1} \quad (7.49)$$

and find the image of the unit circle $|r_2| = 1$ on the z plane. The image of the unit circle on the $-r_2S_{22}$ plane is a circle with its center at the origin and radius $|S_{22}|$. The image on the $(1 - r_2S_{22})$ plane is obtained by shifting the previous image to the right by a unit distance. The image on the $(1 - r_2S_{22})^{-1}$ plane is a circle with its center on the real axis and with intersection points on the real axis at $(1 + |S_{22}|)^{-1}$ and $(1 - |S_{22}|)^{-1}$. On the $B(1 - r_2S_{22})^{-1}$ plane, it becomes a circle with the center at

$$\frac{1}{2} \left(\frac{B}{1 - |S_{22}|} + \frac{B}{1 + |S_{22}|} \right) = \frac{B}{1 - |S_{22}|^2}$$

and radius

$$\frac{1}{2} \left(\frac{|B|}{1 - |S_{22}|} - \frac{|B|}{1 + |S_{22}|} \right) = \frac{|B| |S_{22}|}{1 - |S_{22}|^2}$$

Note that the inside of the unit circle $|r_2| = 1$ corresponds to the inside of this circle since $|S_{22}| < 1$. The complex vector representation of this circle is given by

$$\varrho = \frac{B}{1 - |S_{22}|^2} + \frac{|B| |S_{22}|}{1 - |S_{22}|^2} e^{j\theta}$$

where θ varies from 0 to 2π and ϱ represents the vector drawn from the origin to the moving point on the circle. From this, we see that the unit circle $|r_2| = 1$ is transformed by (7.49) into a circle represented by

$$\begin{aligned} z &= -\frac{A}{S_{22}} + \frac{B}{1 - |S_{22}|^2} + \frac{|B| |S_{22}|}{1 - |S_{22}|^2} e^{j\theta} \\ &= \frac{AS_{22}^* + S_{11}}{1 - |S_{22}|^2} + \frac{|B| |S_{22}|}{1 - |S_{22}|^2} e^{j\theta} \end{aligned} \quad (7.50)$$

If the real part of the input impedance is to remain positive, the magnitude of z has to be smaller than unity regardless of the value of θ , i.e.,

$$\frac{|AS_{22}^* + S_{11}| + |B| |S_{22}|}{1 - |S_{22}|^2} < 1$$

The above inequality is equivalent to

$$S_{11} + \frac{S_{22}^* S_{12} S_{21}}{1 - |S_{22}|^2} < 1 - \frac{|S_{12} S_{21}|}{1 - |S_{22}|^2} \quad (7.51)$$

Let us assume that

$$|S_{12} S_{21}| < 1 - |S_{22}|^2 \quad (7.52)$$

otherwise (7.51) does not hold. Squaring both sides of (7.51), we obtain after a little manipulation

$$2 \operatorname{Re} S_{11} S_{22} S_{12}^* S_{21}^* < 1 - |S_{11}|^2 - |S_{22}|^2 + |S_{11} S_{22}|^2 + |S_{12} S_{21}|^2 - 2 |S_{12} S_{21}|$$

which reduces to

$$|S_{11}|^2 + |S_{22}|^2 + 2 |S_{12} S_{21}| < 1 + |S_{12} S_{21} - S_{11} S_{22}|^2 \quad (7.53)$$

Thus, we see that the real part of the input impedance remains positive if

(7.52) and (7.53) are satisfied. The corresponding conditions for the real part of the output impedance to remain positive are obtained by interchanging the subscripts 1 and 2. However, since (7.53) remains unchanged during this interchange, the conditions for the amplifier to be unconditionally stable are given by

$$\begin{aligned} |S_{12}S_{21}| &< 1 - |S_{11}|^2 \\ |S_{12}S_{21}| &< 1 - |S_{22}|^2 \\ |S_{11}|^2 + |S_{22}|^2 + 2|S_{12}S_{21}| &< 1 + |S_{12}S_{21} - S_{11}S_{22}|^2 \end{aligned} \quad (7.54)$$

Note that (7.54) is a necessary and sufficient condition for unconditional stability.

If an amplifier is unconditionally stable, it can be simultaneously matched to generator and load impedances with positive real parts. For the proof of this statement, we proceed as follows. When the load impedance Z_L is connected to the output port, from (7.46) the reflection coefficient at the input port is given by

$$(Z_i - Z_1^*)/(Z_i + Z_1) = S_{11} + S_{12}r_2S_{21}(1 - r_2S_{22})^{-1}$$

where Z_i is the corresponding input impedance. Let r_1 be the reflection coefficient of Z_g with respect to Z_1^* , i.e.,

$$r_1 = (Z_g - Z_1)/(Z_g + Z_1^*) \quad (7.55)$$

and suppose that r_1^* is equal to $(Z_i - Z_1^*)/(Z_i + Z_1)$, then Z_g becomes the complex conjugate of Z_i indicating that the input port is matched. Thus,

$$r_1^* = S_{11} + S_{12}r_2S_{21}(1 - r_2S_{22})^{-1} \quad (7.56)$$

gives the matching condition at the input port when the load impedance is Z_L . Similarly, when the generator impedance Z_g is connected to the input port, the matching condition at the output port is given by

$$r_2^* = S_{22} + S_{21}r_1S_{12}(1 - r_1S_{11})^{-1} \quad (7.57)$$

For simultaneous matchings, these two equations must both be satisfied. To simplify the following discussion, let us now assume that either one of S_{11} or S_{22} , say S_{22} , is equal to zero. This can be done without loss of generality when we discuss unconditionally stable amplifiers since, for a given Z_1 with a positive real part, Z_2 with a positive real part can always be chosen to satisfy the matching condition $S_{22} = 0$. Under this assumption, eliminating r_2 from (7.56) and (7.57), we have

$$r_1^2 S_{11} + r_1(|S_{12}S_{21}|^2 - |S_{11}|^2 - 1) + S_{11}^* = 0 \quad (7.58)$$

The solutions of this equation are given by

$$r_1 = \frac{1 + |S_{11}|^2 - |S_{12}S_{21}|^2 \pm \{(1 + |S_{11}|^2 - |S_{12}S_{21}|^2)^2 - 4|S_{11}|^2\}^{1/2}}{2S_{11}} \quad (7.59)$$

where $S_{11} \neq 0$ is assumed; otherwise the input is matched to Z_1 , and no discussion is necessary. The expression inside the square root is positive since (7.54) reduces to $1 - |S_{11}| > |S_{12}S_{21}|$ which is equivalent to

$$1 + |S_{11}|^2 - |S_{12}S_{21}|^2 > 2|S_{11}|$$

The product of two solutions of (7.58) is equal to S_{11}^*/S_{11} . Hence, one of the solutions should have a magnitude smaller than one. With this solution, the simultaneous matchings can be achieved keeping the real part of the generator impedance positive. Since the output impedance has a positive real part under this condition, the matched load impedance must also have a positive real part, which completes the proof.

From the above discussion concerning unconditionally stable amplifiers, we see that S_{11} and S_{22} can be set equal to zero from the beginning by choosing proper values for Z_1 and Z_2 , each with a positive real part. This greatly simplifies the discussion, for example, if $S_{11} = S_{22} = 0$, (7.54) reduces to

$$|S_{12}S_{21}| < 1 \quad (7.60)$$

For a given unconditionally stable amplifier and a fixed generator impedance, let us consider the effect of varying the load impedance on the transducer gain. The maximum transducer gain is obtained when the output port is matched since the generator and the amplifier together can be considered as a new generator with an internal impedance having a positive real part. Next, let us consider the case in which the generator impedance is varied, keeping the available power of the generator constant. For a given load impedance, the output power, and hence the transducer gain is maximized when the input port is matched since the load current is proportional to the input current which is maximized when the input port is matched. Now suppose that we simultaneously change the generator and load impedances keeping the output port matched. Suppose also that under this condition, the generator impedance Z_1 gives the highest transducer gain and let Z_2 be the corresponding load impedance, then the output port is guaranteed to be matched. If the input port is not matched, then the same load impedance Z_2 has a corresponding generator impedance Z_1' which gives the matching condition at the input port and hence a higher transducer

gain. With Z_1' , the output port may not be matched. However, there is a load impedance Z_2' which satisfies the matching condition at the output port, and hence gives a still higher transducer gain. This contradicts the assumption that Z_1 gives the highest transducer gain under the matched condition at the output port. We, therefore, conclude that when the generator and load impedances of an unconditionally stable amplifier are allowed to change arbitrarily, provided their real parts are positive, the maximum transducer gain is obtained when both ports are simultaneously matched.

Let us consider the canonical form of an unconditionally stable amplifier. If the amplifier does not contain a positive resistance, the real part of the impedance looking into one port of the amplifier must become negative when the termination at the other port becomes almost purely reactive. Consequently, there must be a positive and a negative resistance in the canonical form of an unconditionally stable amplifier to ensure the impedance has a positive real part. The equivalent circuit is shown in Fig. 7.10. Let Z_1 and Z_2 be the reference impedances at the input and output ports used to define the scattering matrix such that $\text{Re } Z_1 > 0$, $\text{Re } Z_2 > 0$, $S_{11} = 0$, $S_{22} = 0$. Note that we can assume $|S_{21}| > 1$, since $|S_{21}|^2$ is the maximum transducer gain of the amplifier.

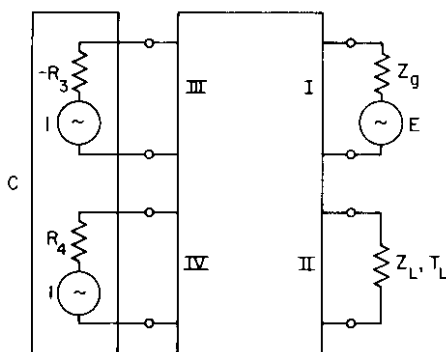


Fig. 7.10. Equivalent circuit of an unconditionally stable amplifier.

Now assume that port IV in Fig. 7.10 can be effectively disconnected from port II by choosing a proper generator impedance, then M_{opt} is realized provided an appropriate load impedance is chosen to satisfy the matching condition at the output port. The existence of such a load impedance with a positive real part is guaranteed by the unconditional stability. Let us now investigate whether or not ports II and IV can be effectively

disconnected as assumed above by connecting some Z_g with a positive real part at port I. When Z_g (with no voltage source in series) is connected to port I,

$$a_1 = r_1 b_1 \quad (7.61)$$

where r_1 is the reflection coefficient of Z_g with respect to Z_1^* as shown in (7.55). Since the voltage sources in Fig. 7.10 are independent of each other and the circuit is linear, the principle of superposition holds, and we can assume that $a_2 = a_3 = 0$ when investigating the effect of a_4 on the output port. Remembering that $S_{11} = 0$, the first two scattering equations then become

$$b_1 = S_{14}a_4, \quad b_2 = S_{21}a_1 + S_{24}a_4$$

Substituting (7.61) in the second equation and solving for b_2 , we obtain

$$b_2 = (S_{21}r_1 S_{14} + S_{24})a_4$$

Therefore, if we choose Z_g such that

$$r_1 = -(S_{24}/S_{21}S_{14}) \quad (7.62)$$

then ports II and IV are effectively disconnected. In order to ensure $\text{Re } Z_g > 0$, we have to show that $|r_1| < 1$. From the lossless condition (7.31), the 11, 22, and 12 components give

$$\begin{aligned} |S_{12}|^2 - 1 + |S_{14}|^2 &= |S_{13}|^2 \\ |S_{21}|^2 - 1 + |S_{24}|^2 &= |S_{23}|^2 \\ |S_{13}S_{23}| &= |S_{14}S_{24}| \end{aligned}$$

respectively. Multiplying the left-hand sides of the first and second equations and equating the result to $|S_{14}S_{24}|^2$, which is equal to $|S_{13}S_{23}|^2$ from the equation, we obtain

$$(|S_{21}|^2 - 1)(|S_{12}|^2 - 1) + |S_{24}|^2(|S_{12}|^2 - 1) + |S_{14}|^2(|S_{21}|^2 - 1) = 0$$

which is equivalent to

$$(1 - |S_{12}|^2) \{1 + |S_{24}|^2(|S_{21}|^2 - 1)^{-1}\} = |S_{14}|^2 \quad (7.63)$$

Since $|S_{21}|^2 > 1$, and from (7.60)

$$|S_{12}|^2 < |S_{21}|^{-2}$$

the substitution of $|S_{12}|^2$ by $|S_{21}|^{-2}$ in (7.63) gives

$$(1 - |S_{21}|^{-2}) \{1 + |S_{24}|^2(|S_{21}|^2 - 1)^{-1}\} < |S_{14}|^2$$

which reduces to

$$|S_{21}|^2 - 1 + |S_{24}|^2 < |S_{14}S_{21}|^2$$

However, since $|S_{21}|^2 > 1$, this shows that

$$|S_{24}|^2 < |S_{14}S_{21}|^2$$

which guarantees $|r_1| < 1$ from (7.62).

From the above discussion, we conclude that M_{opt} can be realized by choosing proper generator and load impedances when the amplifier is unconditionally stable. In other words, neither feedback nor a circulator is necessary to obtain the optimum noise measure. If the generator and load impedances are specified beforehand, we can insert lossless two-port networks at the input and output ports to transform these impedances to appropriate values. The unconditional stability at one frequency does not guarantee the same stability at other frequencies. If adjustable circuits are inserted at the input and output in order to obtain M_{opt} , the system may oscillate at some other frequency. The circuits which transform the generator, and load impedances for M_{opt} must, therefore, be designed carefully in order to maintain negligibly small insertion loss at the desired frequency while at the same time preventing possible oscillations at other frequencies.

7.4 Unilateral Gain

In Section 7.1, it was shown that the values of the resistances in the canonical form are invariant to nonsingular lossless transformations. These invariants are unique in the sense that they and their functions are the only invariants. However, if the transformations are assumed to be reciprocal as well as lossless, we can expect still more invariants. In fact, there is at least one more invariant called the unilateral gain which is defined by

$$U = \frac{|\det(\mathbf{Z} - \mathbf{Z}_r)|}{\det(\mathbf{Z} + \mathbf{Z}^*)} \quad (7.64)$$

In this section, we shall prove the invariance and then rewrite the formula in terms of the scattering matrix in order to see the origin of its name.

Suppose that an n -port network having an impedance matrix $\mathbf{Z}$ is transformed into another n -port network by a nonsingular, lossless, reciprocal $2n$ -port network. The new impedance matrix $\mathbf{Z}'$ is given by (7.13). The lossless condition for the $2n$ -port network is expressed by (7.15). The

reciprocal condition gives

$$\mathbf{Z}_{aa} = \mathbf{Z}_{aat}, \quad \mathbf{Z}_{ab} = \mathbf{Z}_{bat}, \quad \mathbf{Z}_{bb} = \mathbf{Z}_{bbt} \quad (7.65)$$

Combining (7.15) and (7.65), we see that the component matrices are expressible in terms of real matrices $\mathbf{X}_{aa}$, $\mathbf{X}_{ab}$, $\mathbf{X}_{bb}$ as follows:

$$\mathbf{Z}_{aa} = j\mathbf{X}_{aa}, \quad \mathbf{Z}_{ab} = j\mathbf{X}_{ab}, \quad \mathbf{Z}_{bb} = j\mathbf{X}_{bb}$$

where $\mathbf{X}_{aa}$ and $\mathbf{X}_{bb}$ are both symmetrical with respect to the main diagonal, and $\mathbf{X}_{ab} = (\mathbf{X}_{ba})^t$. In terms of these real matrices, (7.13) becomes

$$\mathbf{Z}' = \mathbf{X}_{ba}(\mathbf{Z} + j\mathbf{X}_{aa})^{-1} \mathbf{X}_{ab} + j\mathbf{X}_{bb} \quad (7.66)$$

An inspection of (7.66) shows that the impedance transformation can be decomposed into three different types of transformation:

- (i) $\mathbf{Z}' = \mathbf{Z} + j\mathbf{X}$, where $\mathbf{X}_t = \mathbf{X}$;
- (ii) $\mathbf{Z}' = \mathbf{X}\mathbf{Z}\mathbf{X}_t$, where $\mathbf{X}$ represents an arbitrary real matrix;
- (iii) $\mathbf{Z}' = \mathbf{Z}^{-1}$.

If a certain quantity is invariant to these three transformations, it is also invariant to the transformation given by (7.66). For the proof of the invariance of U , therefore, we have only to show the invariance during these transformations. For transformation (i), we have

$$\begin{aligned} \mathbf{Z}' - \mathbf{Z}'_t &= \mathbf{Z} + j\mathbf{X} - (\mathbf{Z}_t + j\mathbf{X}_t) = \mathbf{Z} - \mathbf{Z}_t \\ \mathbf{Z}' + \mathbf{Z}'^* &= \mathbf{Z} + j\mathbf{X} + (\mathbf{Z}^* - j\mathbf{X}) = \mathbf{Z} + \mathbf{Z}^* \end{aligned}$$

and hence U is invariant. For transformation (ii), we have

$$\begin{aligned} \mathbf{Z}' - \mathbf{Z}'_t &= \mathbf{X}\mathbf{Z}\mathbf{X}_t - (\mathbf{X}\mathbf{Z}\mathbf{X}_t)_t = \mathbf{X}(\mathbf{Z} - \mathbf{Z}_t)\mathbf{X}_t, \\ \mathbf{Z}' + \mathbf{Z}'^* &= \mathbf{X}\mathbf{Z}\mathbf{X}_t + (\mathbf{X}\mathbf{Z}\mathbf{X}_t)^* = \mathbf{X}(\mathbf{Z} + \mathbf{Z}^*)\mathbf{X}_t \end{aligned}$$

Using these relations, U' is calculated to be

$$U' = \frac{|\det(\mathbf{Z}' - \mathbf{Z}'_t)|}{\det(\mathbf{Z}' + \mathbf{Z}'^*)} = \frac{|\det \mathbf{X}|^2 |\det(\mathbf{Z} - \mathbf{Z}_t)|}{(\det \mathbf{X})^2 \det(\mathbf{Z} + \mathbf{Z}^*)} = U$$

and hence U is again invariant. Finally, for transformation (iii) we have

$$\begin{aligned} \mathbf{Z}' - \mathbf{Z}'_t &= \mathbf{Z}^{-1}(\mathbf{Z} - \mathbf{Z}_t)(-\mathbf{Z}_t)^{-1} \\ \mathbf{Z}' + \mathbf{Z}'^* &= \mathbf{Z}^{-1}(\mathbf{Z} + \mathbf{Z}^*)(\mathbf{Z}^*)^{-1} \end{aligned}$$

and hence

$$U' = \frac{|\det(\mathbf{Z}' - \mathbf{Z}_t')|}{\det(\mathbf{Z}' + \mathbf{Z}^*)} = \frac{|\det \mathbf{Z}^{-1}|^2 |\det(\mathbf{Z} - \mathbf{Z}_t)|}{|\det \mathbf{Z}^{-1}|^2 \det(\mathbf{Z} + \mathbf{Z}^*)} = U$$

This completes the proof that U is invariant to nonsingular, lossless, reciprocal transformations.

Let us next consider how to express U in terms of the scattering matrix $\mathbf{S}$. Using (5.79), $\mathbf{Z} - \mathbf{Z}_t$ can be rewritten in terms of $\mathbf{S}$ as follows:

$$\begin{aligned} \mathbf{Z} - \mathbf{Z}_t &= \mathbf{F}^{-1}(\mathbf{I} - \mathbf{S})^{-1} \{(\mathbf{S}\mathbf{G} + \mathbf{G}^+) \mathbf{F}^2(\mathbf{I} - \mathbf{S}_t) \\ &\quad - (\mathbf{I} - \mathbf{S}) \mathbf{F}^2(\mathbf{G}\mathbf{S}_t + \mathbf{G}^+)\} (\mathbf{I} - \mathbf{S}_t)^{-1} \mathbf{F}^{-1} \\ &= \mathbf{F}^{-1}(\mathbf{I} - \mathbf{S})^{-1} \{\mathbf{S}\mathbf{F}^2(\mathbf{G} + \mathbf{G}^+) - \mathbf{F}^2(\mathbf{G} + \mathbf{G}^+) \mathbf{S}_t\} (\mathbf{I} - \mathbf{S}_t)^{-1} \mathbf{F}^{-1} \\ &= \frac{1}{2} \mathbf{F}^{-1}(\mathbf{I} - \mathbf{S})^{-1} (\mathbf{S}\mathbf{P} - \mathbf{P}\mathbf{S}_t) (\mathbf{I} - \mathbf{S}_t)^{-1} \mathbf{F}^{-1} \end{aligned} \quad (7.67)$$

Similarly, $\mathbf{Z} + \mathbf{Z}^*$ becomes

$$\begin{aligned} \mathbf{Z} + \mathbf{Z}^* &= \mathbf{F}^{-1} \{(\mathbf{I} - \mathbf{S})^{-1} (\mathbf{S}\mathbf{G} - \mathbf{G} + \mathbf{G} + \mathbf{G}^+) \\ &\quad + (\mathbf{I} - \mathbf{S}^*)^{-1} (\mathbf{S}^*\mathbf{G}^+ - \mathbf{G}^+ + \mathbf{G}^+ + \mathbf{G})\} \mathbf{F} \\ &= \mathbf{F}^{-1} \{-\mathbf{I} + (\mathbf{I} - \mathbf{S})^{-1} + (\mathbf{I} - \mathbf{S}^*)^{-1}\} (\mathbf{G} + \mathbf{G}^+) \mathbf{F} \\ &= \frac{1}{2} \mathbf{F}^{-1}(\mathbf{I} - \mathbf{S})^{-1} \{- (\mathbf{I} - \mathbf{S})(\mathbf{I} - \mathbf{S}^*) \\ &\quad + (\mathbf{I} - \mathbf{S}^*) + (\mathbf{I} - \mathbf{S})\} (\mathbf{I} - \mathbf{S}^*)^{-1} \mathbf{P}\mathbf{F}^{-1} \\ &= \frac{1}{2} \mathbf{F}^{-1}(\mathbf{I} - \mathbf{S})^{-1} (\mathbf{I} - \mathbf{S}\mathbf{S}^*) (\mathbf{I} - \mathbf{S}^*)^{-1} \mathbf{P}\mathbf{F}^{-1} \end{aligned} \quad (7.68)$$

Substituting (7.67) and (7.68) into (7.64), we obtain the desired expression,

$$U = \frac{|\det(\mathbf{S}\mathbf{P} - \mathbf{P}\mathbf{S}_t)|}{\det \mathbf{P} \det(\mathbf{I} - \mathbf{S}\mathbf{S}^*)} \quad (7.69)$$

Suppose that a two-port network with unilateral gain U is imbedded in a nonsingular, lossless, reciprocal circuit and that S_{12} of the resultant circuit becomes zero; i.e., the circuit is unilateralized. The U of the unilateralized circuit should remain the same as before because of the invariant property. With $S_{12} = 0$, the generator and load impedances do not affect the output and input impedances, respectively, thereby permitting the input and output ports to be matched independently of each other. This means that it is always possible to assume that $S_{11} = 0$ and $S_{22} = 0$ for the unilateralized circuit by using proper reference impedances. Equation (7.69) then becomes

$$U = p_1 p_2 |S_{21}|^2$$

When $p_1 = p_2 = 1$, U represents the maximum transducer gain of the unilateral circuit which explains why U is called the unilateral gain.

A method for unilateralizing a two-port network is illustrated in Fig. 7.11. The box between ports I and II represents the given two-port network to be unilateralized. First, a reactance jX is connected in series with port II to bring V_2' into phase with V_1 when $I_1 = 0$. Next the ideal transformer is adjusted to cancel V_1 so that V_1' becomes zero. Z'_{12} of the transformed

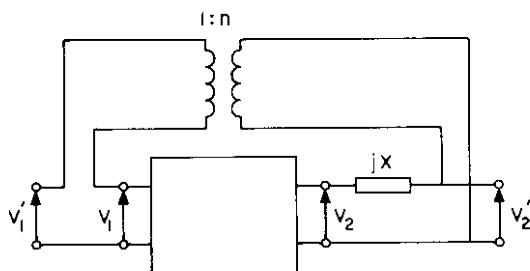


Fig. 7.11. Unilateralization of a two-port network.

circuit then becomes zero by definition, and the corresponding S_{12} vanishes from (5.78), i.e., the circuit is unilateralized.

7.5 Tunnel Diode Amplifiers

If we measure the terminal characteristic of a tunnel diode, we shall find a certain range of bias voltage in which the current decreases as we increase the voltage as shown in Fig. 7.12. Let us fix the bias voltage V_0 somewhere

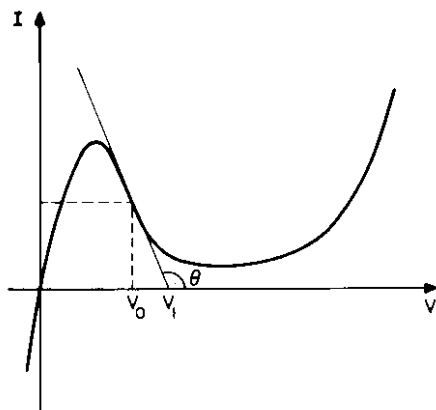


Fig. 7.12. I - V characteristic of a tunnel diode.

in this range and superpose a small microwave signal. Then, when the alternating voltage swings in the positive direction, the corresponding alternating current swings in the negative direction and the device exhibits a negative resistance characteristic. The magnitude of the negative resistance is given by the inverse of the slope of the I - V characteristic at the bias point; i.e., $\cot \theta = -R$. Once a negative resistance becomes available, an amplifier can be built as we discussed at the end of Section 5.8. Before we discuss the amplifier, let us study the origin of the negative resistance characteristic and then construct a simple equivalent circuit for a tunnel diode. To do so, we must first give an elementary discussion of semiconductors.

As is well known, light is one form of electromagnetic energy that propagates as a wave. However, when it strikes a photoelectric surface, it behaves as though it were a particle. To express this corpuscular nature of light, the name photon is ordinarily used. The photon is considered to be associated with the light wave. The probability that a photon is found at a certain place is proportional to the square of the wave amplitude at that point. In much the same way, an electron is also associated with a wave. Inside an atom, a standing wave is formed for each electron around the nucleus, and the probability of finding an electron at a certain point is proportional to the square of the amplitude of the corresponding standing wave at that point. Only certain functions are acceptable for representing such standing waves since each has to be a single-valued continuous function of position while its integral over the whole space must be finite. The electron associated with a certain standing wave has a certain energy determined by the corresponding wave function. Therefore, in a steady state, electrons inside an atom can have only discrete energies. This situation is similar to that of a resonant cavity in which only certain field patterns are allowed to exist, and each has its own resonant frequency.

If we consider the electron as a particle, it has a certain potential energy due to the attractive force from the positive charge of the nucleus, and kinetic energy due to its motion inside the atom. The graph of potential energy as a function of the distance from the nucleus should look similar to a morning-glory flower whose cross section is shown in Fig. 7.13. Suppose that the electron has total energy E_1 and that a horizontal line is drawn at $E = E_1$. The region between the two intersections with the potential curve represents the space within which the electron can exist as a classical particle. At each intersection, all the kinetic energy is replaced by potential energy, and hence the particle with energy E_1 is not permitted to move beyond this point. It could be imagined that there is a fictitious "potential wall" along

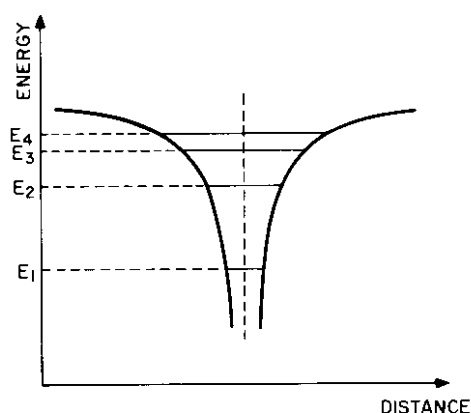


Fig. 7.13. Potential energy as a function of distance in an atom.

the potential curve. However, the standing wave associated with the electron does not necessarily vanish beyond the potential wall since around the atom there is nothing which behaves like a perfect conductor to reflect the wave completely. This means that there is a finite probability of finding the electron beyond the potential wall; therefore, the horizontal solid line representing the space within which the electron is allowed to exist should be considered to extend beyond the potential wall, tapering off quickly with increasing distance.

There is a limitation to the number of electrons which can belong to each wave function, i.e., which are associated with the same standing wave pattern. No more than two electrons can belong to one wave function, and they must have the opposite spin orientations if they belong to one wave function. When two electrons belong to one wave function, an additional electron has to belong to another wave function, generally with different energy. Furthermore, the first two electrons usually belong to the wave function with the lowest energy, and the next two electrons to the wave function with the next lowest energy, and so forth, until all the electrons in the atom belong to certain functions with the lowest possible energy.

So far we have discussed electrons in a single atom. Next, we discuss many atoms which come close together to form a crystal. As atoms approach one another, the electron wave functions of different atoms start to interfere with each other, and each energy level splits into many different levels. This is similar to the split of the propagation constant of two identical waveguides coupled together. When the crystal is formed, some of the wave functions

extending over large distances from the nuclei interfere with each other so strongly that they lose their identity. In other words, it is no longer clear to which atom they belong, and they spread over the entire crystal. Their energy levels are crowded together and occupy some form of band as shown in Fig. 7.14. Each horizontal line roughly shows the range where the electron belonging to the corresponding wave function can exist. Each wave function

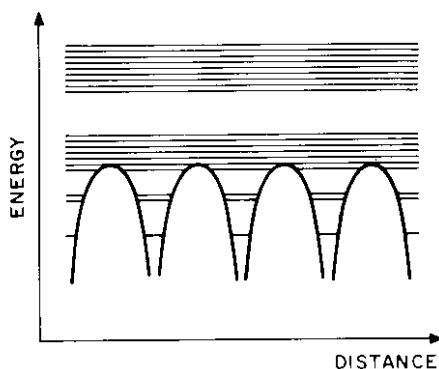


Fig. 7.14. Energy bands in crystal.

can accommodate no more than two electrons as before; and, under this condition, electrons try to belong to the wave function with the lowest possible energy, or equivalently, to be in the lowest possible energy state. Therefore, all the wave functions up to a certain level are filled with electrons, and wave functions above that level are left empty. However, if the atoms vibrate strongly due to heat, some of the electrons obtain energy from the vibration and jump up to the higher energy states which were originally empty. The probability $f(E)$ that electrons occupy a certain energy state therefore appears as shown in Fig. 7.15. At absolute zero temperature, the probability is unity up to a certain level and zero above that level as shown in Fig. 7.15(a). As the temperature increases, the probability that electrons take high energy states increases, and the probability for low energy states correspondingly decreases as shown in Fig. 7.15(b). This is the Fermi-Dirac distribution law. The horizontal dotted line in Fig. 7.15(b) shows the Fermi levels where the probability becomes one half.

Various phenomena of electric conduction can be explained using the energy bands described above and the Fermi-Dirac distribution law. Ordinary metals have an energy band which is only half filled by electrons with the

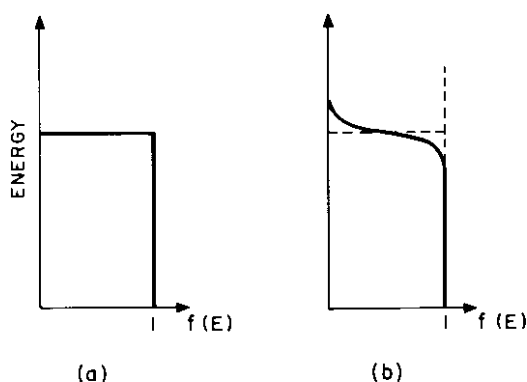


Fig. 7.15. Probability $f(E)$ of electron occupancy of available states: (a) $T = 0^\circ\text{K}$; (b) At a higher temperature.

upper part left empty. When an electric field is applied, the electrons in this band move in the opposite direction associated with the wave packets made up of several empty wave functions just above the original states. A large field is not necessary to start net electron motion or conduction; hence, metals are good conductors. The half-filled energy band is called the conduction band.

For insulators, one energy band called the valence band is completely filled with electrons and the conduction band above it is empty. The energy gap between them is so large that the electrons in the valence band can hardly be excited into the conduction band. Since there is no room for electrons to move in the valence band and no electrons exist in the conduction band, electric current cannot be induced by an ordinary electric field, which is the basic property of insulators.

For intrinsic semiconductors, the energy gap between the valence and conduction band is small, and some electrons from the valence band can be excited into the conduction band by heat leaving behind some room for electrons in the valence band to move. The vacancy each electron leaves behind in the valence band, when it is excited into an upper level, is called a hole. When electrons in the valence band move about using these vacancies, the holes may be considered to move in the opposite direction. Holes, therefore, behave like particles each with a positive electric charge e , where $-e$ is the electron charge. When an electric field is applied, electrons and holes move in the conduction band and valence band, respectively, and contribute to the conduction current. However, the number of electrons

and holes available for this net motion is small, and the conductivity is small compared to metals, which is the reason for the name semiconductor.

For impurity semiconductors, Fig. 7.14 is modified by the presence of impurity atoms as shown in Fig. 7.16. When the energy level of an empty wave function of the impurity atom is located just above the valence band as shown in Fig. 7.16(a), the semiconductor is said to be a p -type. On the other hand, when the energy level of a filled wave function of the impurity atom is located just below the conduction band as shown in Fig. 7.16(b), the semiconductor is said to be an n -type. Electrons in the valence band in a p -type semiconductor can be excited by heat to jump into the empty level

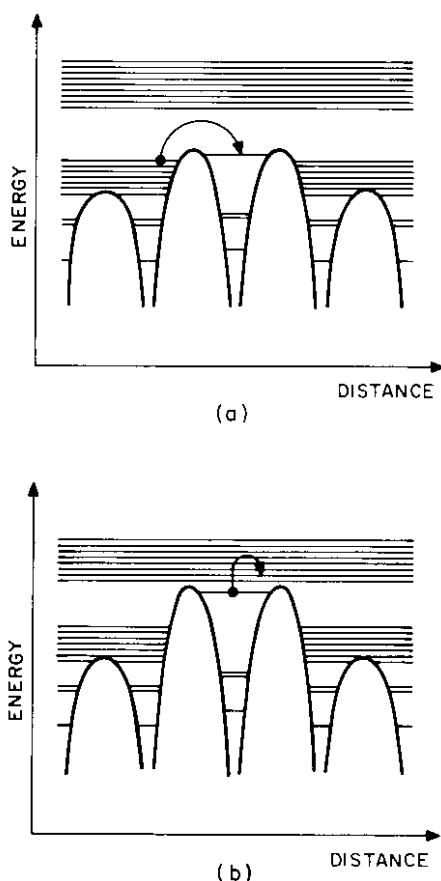


Fig. 7.16. Energy bands and impurity levels in impurity semiconductors: (a) p -type semiconductor; (b) n -type semiconductor.

of the impurity atoms creating holes in the valence band. Similarly, the electrons of the impurity atoms in an n -type semiconductor can jump into the conduction band ready to participate in electric conduction. In either case, conduction current takes place when an electric field is applied. If we indicate the band edges and the energy levels of the impurity atoms, neglecting the potential walls of atoms as well as the detailed structure in the atoms, Figs. 7.16(a) and (b) become Figs. 7.17(a) and (b), respectively. The Fermi level of each case is indicated by the dotted line. At room temperature, some of the impurity atoms in the p -type semiconductor receive electrons from the valence band and are negatively charged or ionized. On the other hand, some of the impurity atoms in the n -type semiconductor give up electrons to the the conduction band and are positively ionized, as illustrated in Fig. 7.17. Although the impurity atoms are charged, the total number of

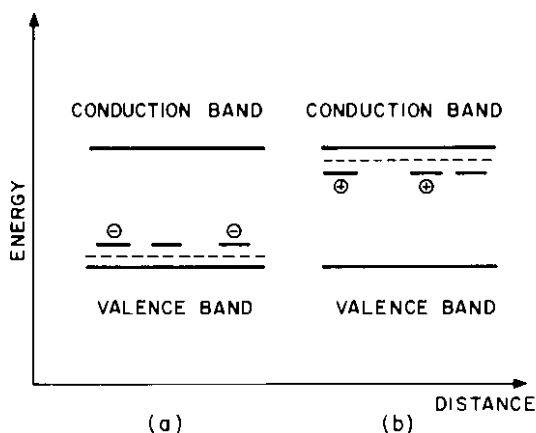


Fig. 7.17. Energy band structures of impurity semiconductors: (a) p -type semiconductor; (b) n -type semiconductor.

electrons in the crystal remains the same, and hence the crystal as a whole is neutral.

Now, suppose that the p -type and n -type semiconductors come into contact. If the Fermi level of the n -side is higher than that of the p -side, electrons move from the n -side into the p -side where more empty states are available at the same energy level; the p -side then becomes negatively charged and the n -side positively charged. Since the vertical axis of our diagram indicates the energy of electrons, the potential is lower towards the top. Thus, the negatively charged p -side rises as a whole and the n -side goes down, as shown

in Fig. 7.18, until the Fermi levels on both sides coincide. Under this condition, the probabilities that available states are occupied by electrons are the same on both sides at each energy level, and no further transfer of electrons takes place. Electrons near the junction in the n -region are transferred into the p -region where they recombine with holes and disappear.

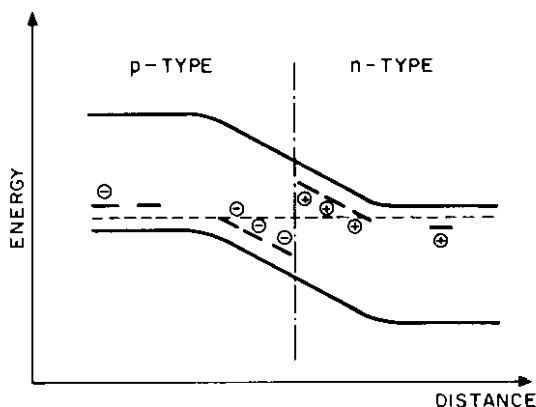


Fig. 7.18. Energy band structure at the p - n junction.

Thus, free charge carriers near the junction are depleted, and the charged impurity atoms build up a strong field which provides the necessary potential difference to make the Fermi levels on both sides coincide.

If we connect a battery to make the p -side negative and the n -side positive, the p -region rises in Fig. 7.18, the n -region goes down, and at the same time the depleted region becomes thicker to accommodate the externally applied potential difference. Both positive and negative charges in the depleted region increase as the region becomes thicker, and since the charges increase as we increase the applied voltage, the junction behaves like a capacitance. This capacitance is called the depletion layer capacitance.

If we connect the battery in the opposite direction, the p -region goes down and the n -region rises, in Fig. 7.18. Electrons then move from the n -region into the p -region where empty states are available abundantly at the same energy level. Since electrons move from the n - to p -sides, net current flows from the p - to n -sides. This is called the forward direction of the junction, and the forward current increases as we increase the applied voltage. When the current flows, excess electrons (holes) in the p -region

(*n*-region) build up a net negative (positive) charge which increases with increasing current, and hence with increasing voltage applied to the junction. This stored charge introduces another capacitance effect called the diffusion capacitance. The depletion layer and diffusion capacitances combined are called the junction capacitance.

Suppose that the number of impurity atoms increases so that one out of every 1000 atoms in the crystal, or equivalently, one out of every ten atoms in one direction is an impurity, then the wave functions of the impurity atoms begin to interfere with each other, and their energy level splits thus forming a band structure, as before. This energy band is connected to the bottom of the conduction band in the case of *n*-type semiconductors and to the top of the valence band in the case of *p*-type semiconductors. The resultant conduction or valence band is now only partially occupied by electrons and behaves like the conduction band in metals. The conductivity of such semiconductors is therefore high. Furthermore, the depletion layer at the junction between two such semiconductors becomes very thin since, with the high density of ionized impurity atoms, only a thin layer is required to build up the potential necessary for the Fermi levels on both sides to coincide. Under this condition, the wave functions on the opposite sides of the depletion layer interfere with each other. The interference may be weak; nevertheless, just as we saw in the discussion of coupled modes, the amplitude of one wave is transferred to the other through the interference, which means the particle in one region is transferred to the other through the interference. Even when no voltage is applied to the junction, electrons are constantly moving from the *p*- to *n*-region and back due to the wave interference. However, since the number of electrons moving in one direction is equal to the number moving in the opposite direction, no net current is observed.

Let us calculate the number of such electrons moving back and forth through the junction. To do so, let $q(E)$ be the number of wave functions in a unit energy interval, counting twice each function for electrons with opposite spin orientations. Let $f(E)$ be the probability of available wave functions being occupied by electrons at energy level E . Subscripts *n* and *p* will be used to express the quantities in the *n*- and *p*-regions. We first consider the transfer from the *n*- to the *p*-region. The number of electrons in the *n*-region in a small energy interval ΔE is given by $q_n(E) \Delta E f_n(E)$, the number of available states times the probability of occupancy. The number of available empty states in the *p*-region is given by $q_p(E) \Delta E \{1 - f_p(E)\}$. Let a be the probability that an electron on one side is transferred to an

empty state on the other side. Then the number of electrons moving from the n - to p -region in ΔE is given by

$$a \varrho_n(E) \Delta E f_n(E) \varrho_p(E) \Delta E \{1 - f_p(E)\}$$

The transfer with a large energy change is unlikely since it corresponds to the interference between two waves with quite different propagation constants. If we neglect the effect of such a transfer, the total number of electrons moving from the n - to the p -region is given by

$$\sum_E (a \Delta E) \varrho_n(E) f_n(E) \varrho_p(E) \{1 - f_p(E)\} \Delta E$$

Similarly, the total number of electrons moving from the p - to the n -region is given by

$$\sum_E (a \Delta E) \varrho_p(E) f_p(E) \varrho_n(E) \{1 - f_n(E)\} \Delta E$$

Therefore, remembering that the electron charge is negative, the net current from the p - to the n -side must be proportional to

$$\sum_E (a \Delta E) \varrho_n(E) \varrho_p(E) \{f_n(E) - f_p(E)\} \Delta E$$

When the Fermi levels on both sides coincide, $f_n(E)$ is equal to $f_p(E)$ for each E and the net current vanishes. If a small forward bias is applied and the p -side potential is raised to lower the Fermi level, $f_n(E) - f_p(E)$ no longer vanishes in the vicinity of the Fermi levels and some net current results. At the edges of the energy bands $\varrho_p(E)$ and $\varrho_n(E)$ generally look like Figs. 7.19(a) and (b), respectively, in which the Fermi levels are shown by

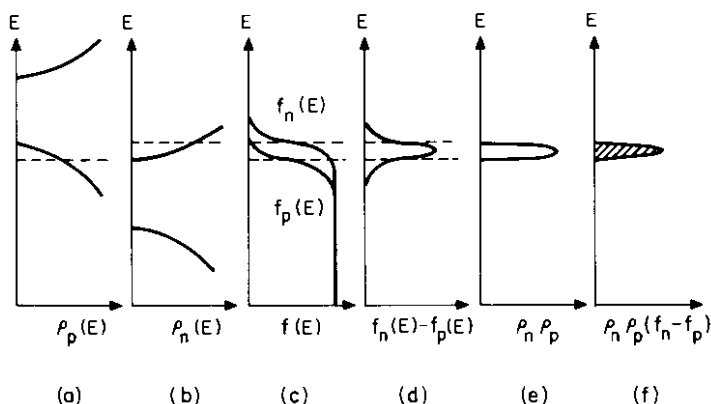


Fig. 7.19. Explanation for the origin of tunnel current.

the dotted lines. Figures 7.19(c), (d), and (e) show the relation between $f_n(E)$ and $f_p(E)$, $f_n(E) - f_p(E)$, and $\varrho_n(E) \varrho_p(E)$, respectively. From these figures, we see that $\varrho_n(e) \varrho_p(e) \{f_n(E) - f_p(E)\}$ should look like Fig. 7.19(f). Since this has nonzero values only over a narrow range E , the transfer probability a can be considered to be a constant over this range, and hence the net current is proportional to the area inside the curve. If the applied voltage is increased further, the range over which $f_n(E) - f_p(E)$ does not vanish increases. On the other hand, the range over which $\varrho_n(E) \varrho_p(E)$ has nonzero values decreases. When the overlap of two energy bands disappears, the area under the curve $\varrho_n(E) \varrho_p(E) \{f_n(E) - f_p(E)\}$ vanishes and the net current disappears. When reverse voltage is applied to the junction, $\{f_n(E) - f_p(E)\}$ becomes negative, and the range for nonzero $\varrho_n(e) \varrho_p(E)$ increases rapidly with increasing reverse voltage. Consequently, the net current due to the wave interference through the junction varies with the applied voltage as shown by the solid line in Fig. 7.20. The dotted line in Fig. 7.20 shows the current through the junction due to the mechanism

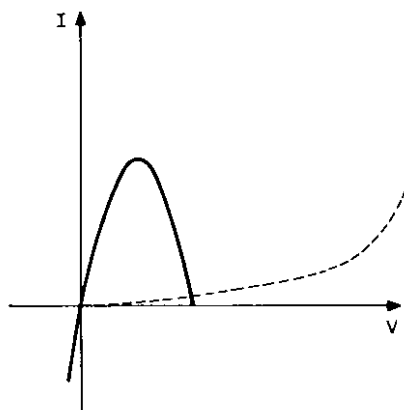


Fig. 7.20. Tunnel current versus applied voltage at the junction.

explained in connection with Fig. 7.18. If we add these two current components, we obtain the I - V characteristic shown in Fig. 7.12.

Classical particles cannot exist behind a potential wall; however, due to their wave nature, electrons can exist behind the wall where classical particles are forbidden. This effect is called the tunnel effect since electrons can penetrate the potential wall as if there were a tunnel through it. The

tunnel diode is named after this effect since it utilizes a similar wave nature of electrons.

In addition to the negative resistance discussed above, the junction of a tunnel diode has capacitance as do other types of junctions. Furthermore, the junction current flows through the semiconductor regions which act as a series resistor. Consequently, an equivalent circuit for the diode becomes the parallel connection of the negative resistance and the junction capacitance in series with the small positive resistance r . In addition, there might be a small series inductance due to the lead wire or whisker connected to the diode wafer inside the capsulation. However, it can be included in the external lossless circuit, and hence it is not considered here.

Let us investigate the possible noise sources in the diode. The small resistance r representing the semiconductor body is an ordinary resistance and has an available noise power kTB when its temperature is $T^\circ\text{K}$. In the series representation, the mean square value of this noise voltage is therefore given by

$$\langle e^2 \rangle = 4kTBr \quad (7.70)$$

In addition, we must consider the shot noise introduced when current flows through the junction. Suppose electrons traverse a certain gap which is not conductive. As each electron approaches the anode an electric charge is induced in the anode which is neutralized by the charge of the electron when it arrives. Hence, the current flowing into the anode consists of a large number of short pulses as illustrated in Fig. 7.21. The average of this current gives the dc component I_0 through the gap. At the same time, a noise component is produced since the position of each pulse is entirely random; this noise is large when I_0 is large. A statistical discussion, which we shall not present here (see Problem 9.5) shows that when the pulse width is

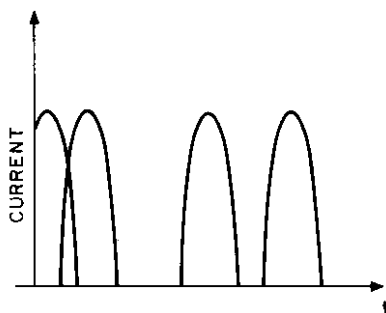


Fig. 7.21. Explanation of shot noise.

sufficiently narrow compared to one cycle of the frequency of interest, the mean square value of the noise current is given by

$$\langle i_n^2 \rangle = 2eI_0B \quad (7.71)$$

This is called the shot noise. The depletion layer at the junction is not conductive, and it produces the shot noise given by (7.71). For I_0 , we have to add the currents due to electrons flowing from the p - to the n -region in addition to those flowing in the reverse direction. Consequently, I_0 is generally larger than the terminal current I which is given by the I - V characteristic of the diode. For most diodes, however, I_0 is only 20–30% larger than I at ordinary operating points, and, for some good diodes, the discrepancy is considerably less. If we include these noise sources in the equivalent circuit of the diode, we obtain Fig. 7.22.

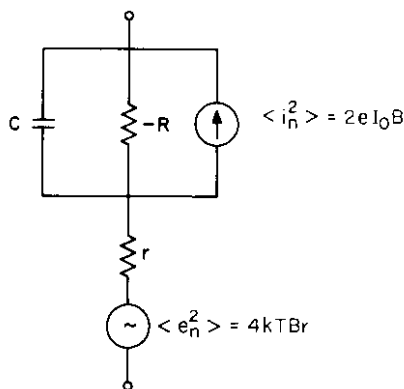


Fig. 7.22. Equivalent circuit of a tunnel diode.

Amplification is possible only when the real part of the impedance looking from the terminals is negative. The impedance Z is given by

$$Z = r + \frac{1}{-(1/R) + j\omega C} = r - \frac{R}{1 + (\omega CR)^2} - j \frac{\omega CR^2}{1 + (\omega CR)^2} \quad (7.72)$$

In order to make the real part negative, ω must be lower than

$$\omega_c = (CR)^{-1} \{(R - r)/r\}^{1/2} \quad (7.73)$$

where ω_c is called the cutoff angular frequency beyond which the diode becomes a passive element. If $R \gg r$, as is usually the case, (7.73) reduces to

$$\omega_c \simeq C^{-1}(Rr)^{-1/2} \quad (R \gg r) \quad (7.74)$$

In order to estimate the optimum noise measure M_{opt} obtainable with a tunnel diode, let us next calculate the open-circuited noise voltage at the terminals. Let v_n represent the noise voltage due to the shot noise, then we have

$$v_n = \frac{i_n}{-(1/R) + j\omega C}$$

and hence

$$\langle v_n^2 \rangle = \frac{R^2}{1 + (\omega CR)^2} \langle i_n^2 \rangle$$

The noise voltage due to the series resistance r and the shot noise voltage are independent of each other. The mean square value of the total noise voltage $\langle V_n^2 \rangle$ is, therefore, given by

$$\langle V_n^2 \rangle = \langle e_n^2 \rangle + \langle v_n^2 \rangle = 4kTBr + \frac{R^2}{1 + (\omega CR)^2} 2eI_0B$$

The exchangeable noise power P_e is equal to $\langle V_n^2 \rangle$ divided by four times the real part of Z in (7.72), and from the discussion given in Section 7.2, M_{opt} becomes

$$\begin{aligned} M_{\text{opt}} &= \frac{-P_e}{kT_i B} = -\frac{\langle V_n^2 \rangle}{4 \operatorname{Re} Z} \frac{1}{kT_i B} \\ &= \left\{ \frac{r}{R-r} + \frac{R}{R-r} \frac{(\omega/\omega_c)^2}{1 - (\omega/\omega_c)^2} \right\} \frac{T}{T_i} + \frac{R}{R-r} \frac{I_0 R}{(2kT_i/e)} \frac{1}{1 - (\omega/\omega_c)^2} \end{aligned} \quad (7.75)$$

which reduces to

$$M_{\text{opt}} \simeq \frac{(\omega/\omega_c)^2}{1 - (\omega/\omega_c)^2} \frac{T}{T_i} + \frac{I_0 R}{(2kT_i/e)} \frac{1}{1 - (\omega/\omega_c)^2} \quad (7.76)$$

when $R \gg r$.

If we draw a straight line tangential to the I - V characteristic at the operating point as shown in Fig. 7.12 and let V_1 be the intersection with the V -axis and V_0 be the bias voltage for the operating point, then the difference $V_1 - V_0$ gives IR . Since I_0 is generally 20–30% larger than I , $I_0 R$ will be obtained by increasing the value of IR by the same percentage. In addition, $2kT_i/e$ is approximately equal to 0.05 V when $T_i = 290^\circ \text{K}$. If the diode is at room temperature, we may assume that $T \simeq T_i$. Substituting these values in (7.76), M_{opt} can be calculated as shown in Fig. 7.23. The abscissa shows ω/ω_c and the ordinate M_{opt} times T_i . [$F_\infty = 10 \log_{10}(M+1)$; $T_i = 290^\circ \text{K}$]. From the curve corresponding to an appropriate value of $I_0 R$, the best M_{opt}

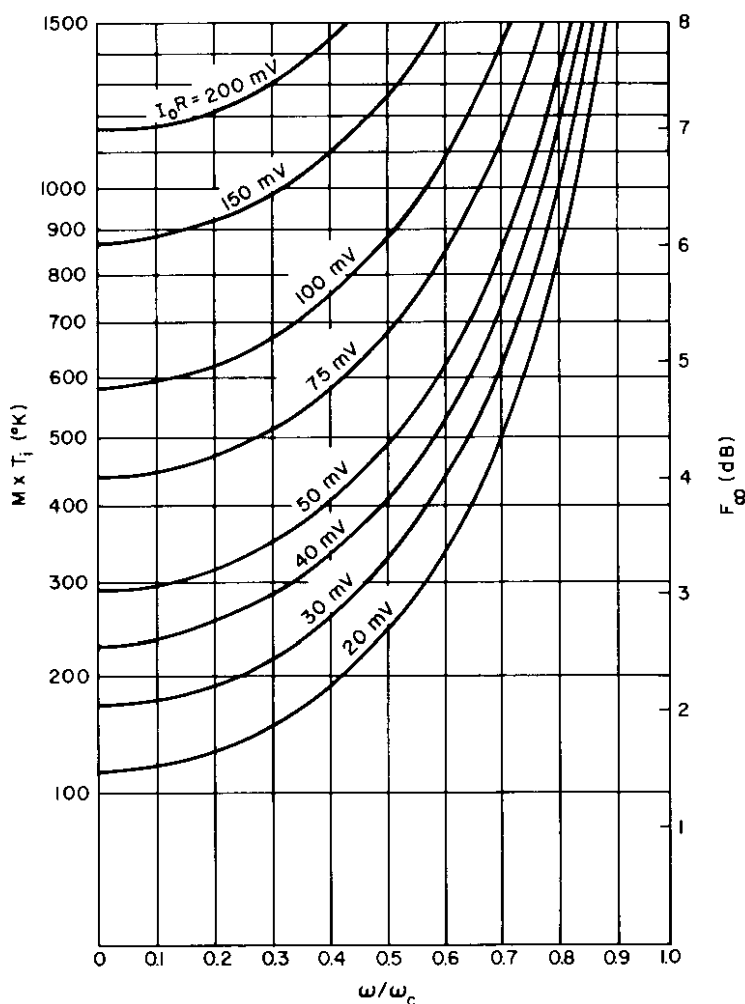
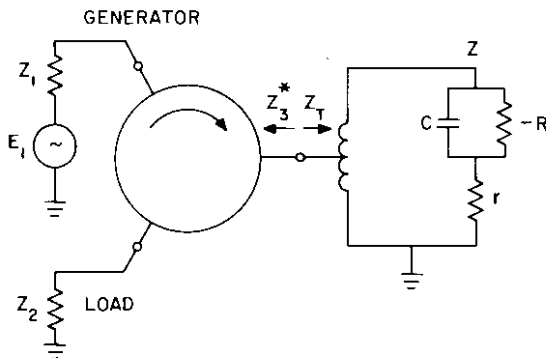


Fig. 7.23. $M_{\text{opt}} \times T_i$ versus ω/ω_c with $I_0 R$ as a parameter.

obtainable with a given tunnel diode can be readily estimated. To realize M_{opt} , the diode impedance is transformed to an appropriate value Z_T through a lossless circuit and is then connected to the generator and load through a lossless circulator as shown in Fig. 7.24. If the reference impedance of the circulator at port III is Z_3 and Z_T is connected to this port, the gain is given by

$$G = |(Z_T - Z_3)/(Z_T + Z_3^*)|^2$$

Fig. 7.24. Realization of M_{opt} .

as we discussed in Section 5.8. If a larger gain is desired, the circuit must be adjusted to make $-Z_T$ close to Z_3^* without exciting possible oscillations.

In practice, the circulator as well as the impedance transforming circuit has some insertion loss; hence, the noise measure realizable with a given tunnel diode is always worse than the value calculated as above. The circulator insertion loss is usually of the order of 0.1–0.3 dB, and the noise temperature ratio T_{op}/T_i of the amplifier in dB increases by approximately the same amount provided the circulator temperature is close to T_i . On the other hand, the transducer gain decreases by twice this amount. Taking these factors into account, the noise measure obtainable in practice can be estimated from (7.37).

7.6 Manley-Rowe Relations

A circuit element which destroys the linear relationship between the terminal voltage and current is described as being nonlinear. When a nonlinear element is lossless, it is called a nonlinear reactance. Suppose voltage at angular frequencies ω_s and ω_p are applied to a nonlinear reactance, then we expect that harmonics at frequencies expressible in the form $|m\omega_s + n\omega_p|$ may appear, where m and n are positive or negative integers. Let us assume that no other frequency components exist and that the two voltage sources are independent. Then we have

$$\sum_{m=0}^{\infty} \sum_{n=-\infty}^{\infty} \frac{mP_{m,n}}{m\omega_s + n\omega_p} = 0 \quad (7.77)$$

$$\sum_{m=-\infty}^{\infty} \sum_{n=0}^{\infty} \frac{nP_{m,n}}{m\omega_s + n\omega_p} = 0 \quad (7.78)$$

where $P_{m,n}$ represents the net power into the element at $|m\omega_s + n\omega_p|$. These two relations are called the Manley-Rowe relations, and we shall derive them for a nonlinear capacitance using the relation

$$v = f(q) \quad (7.79)$$

where v is the terminal voltage, q is the stored charge, and $f(q)$ represents a single-valued function of q . In the derivation for a nonlinear inductance case, we have only to replace (7.79) by $i = f(\varphi)$, where i is the terminal current and φ is the magnetic flux through the inductance.

By hypothesis, the voltage has nonzero components at $|m\omega_s + n\omega_p|$ only, i.e.,

$$\begin{aligned} v &= \sum_{m=0}^{\infty} \sum_{n=-\infty}^{\infty} A_{m,n} \cos(m\omega_s + n\omega_p)t + B_{m,n} \sin(m\omega_s + n\omega_p)t \\ &= \sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} (V_{m,n}/\sqrt{2}) e^{j(m\omega_s + n\omega_p)t} \end{aligned} \quad (7.80)$$

where

$$V_{m,n} = (A_{m,n} - jB_{m,n})/\sqrt{2}, \quad V_{-m,-n} = (A_{m,n} + jB_{m,n})/\sqrt{2}, \quad V_{0,0} = \sqrt{2}A_{0,0}$$

The factor $\sqrt{2}$ in (7.80) is introduced so that $V_{m,n}$ can be interpreted as the effective value of the voltage at $m\omega_s + n\omega_p$. Note that

$$V_{m,n} = V_{-m,-n}^*$$

Similarly, the stored charge can be written in the form

$$q = \sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} (Q_{m,n}/\sqrt{2}) e^{j(m\omega_s + n\omega_p)t} \quad (7.81)$$

where the $Q_{m,n}$'s are expansion coefficients and

$$Q_{m,n} = Q_{-m,-n}^*$$

Let us assume that the infinite series

$$\sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} j(m\omega_s + n\omega_p) Q_{m,n} e^{j(m\omega_s + n\omega_p)t}$$

converges uniformly, then the current can be expressed as:

$$i = (dq/dt) = \sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} (I_{m,n}/\sqrt{2}) e^{j(m\omega_s + n\omega_p)t}$$

where

$$I_{m,n} = j(m\omega_s + n\omega_p) Q_{m,n} \quad (7.82)$$

The average power flowing into the nonlinear capacitance at $|m\omega_s + n\omega_p|$ is given by

$$P_{m,n} = \text{Re} \{V_{m,n} I_{m,n}^*\} = \frac{1}{2} (V_{m,n} I_{m,n}^* + V_{m,n}^* I_{m,n}) \quad (7.83)$$

From the lossless condition, the total power into the capacitance must be equal to zero;

$$\sum_{m,n} P_{m,n} = 0 \quad (7.84)$$

where the summation is taken to include the power corresponding to each frequency only once. To facilitate the following calculation, let us rewrite (7.83) as follows:

$$P_{m,n} = W_{m,n} + W_{-m,-n} \quad (7.85)$$

where

$$\begin{aligned} W_{m,n} &= \frac{1}{4} (V_{m,n} I_{m,n}^* + V_{m,n}^* I_{m,n}) \\ &= \frac{1}{4} (m\omega_s + n\omega_p) (-jV_{m,n} Q_{m,n}^* + jV_{m,n}^* Q_{m,n}). \end{aligned} \quad (7.86)$$

Then, (7.84) can be rewritten in the form

$$\sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} W_{m,n} = 0$$

or equivalently,

$$\omega_s \sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} \frac{mW_{m,n}}{m\omega_s + n\omega_p} + \omega_p \sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} \frac{nW_{m,n}}{m\omega_s + n\omega_p} = 0 \quad (7.87)$$

Since (7.79) holds regardless of the values of ω_s and ω_p , we see from (7.80) and (7.81) that each $V_{m,n}$ can be considered as a function of the $Q_{m,n}$'s only. Then, (7.86) shows that $W_{m,n}/(m\omega_s + n\omega_p)$ is a function of the $Q_{m,n}$'s only and is independent of ω_s and ω_p . Suppose we vary ω_s and ω_p independently while keeping the $Q_{m,n}$'s the same (the external circuit condition will have to be changed as ω_s and ω_p change in order to accomplish this), then from (7.87) we see that

$$\sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} \frac{mW_{m,n}}{m\omega_s + n\omega_p} = 0 \quad (7.88)$$

$$\sum_{m=-\infty}^{\infty} \sum_{n=-\infty}^{\infty} \frac{nW_{m,n}}{m\omega_s + n\omega_p} = 0 \quad (7.89)$$

Otherwise, (7.87) is not satisfied during the variation. Equation (7.88) is

equivalent to

$$\sum_{m=0}^{\infty} \sum_{n=-\infty}^{\infty} \frac{mW_{m,n}}{m\omega_s + n\omega_p} + \sum_{m=-\infty}^0 \sum_{n=-\infty}^{\infty} \frac{mW_{m,n}}{m\omega_s + n\omega_p} = 0 \quad (7.90)$$

Changing the signs of m, n and using the relation $W_{m,n} = W_{-m,-n}$ which can easily be derived from (7.86), the second term on the left-hand side becomes

$$\sum_{m=-\infty}^0 \sum_{n=-\infty}^{\infty} \frac{mW_{m,n}}{m\omega_s + n\omega_p} = \sum_{m=0}^{\infty} \sum_{n=-\infty}^{\infty} \frac{mW_{m,n}}{m\omega_s + n\omega_p}$$

Substituting this into (7.90), we obtain (7.77). Similarly, from (7.89) we obtain (7.78). This completes the derivation of the Manley-Rowe relations.

Let us consider some simple applications of the Manley-Rowe relations. Suppose that a large voltage at ω_p is applied to a nonlinear reactance, for example, to the junction capacitance of a diode, and a small signal voltage at ω_s is superimposed. The large voltage is called the pump voltage, and ω_p the pump (angular) frequency. If the nonlinear reactance is reactively terminated at all frequencies except ω_p , ω_s , and $\omega_u = \omega_p + \omega_s$, no net power flow takes place except at these three frequencies. Rewriting $P_{0,1}$, $P_{1,0}$, and $P_{1,1}$ as P_p , P_s , and P_u , respectively, the Manley-Rowe relations become

$$(P_s/\omega_s) + (P_u/\omega_u) = 0 \quad (7.91)$$

$$(P_p/\omega_p) + (P_u/\omega_u) = 0 \quad (7.92)$$

From (7.91), we have

$$P_u = -P_s(\omega_u/\omega_s)$$

This equation shows that P_u is negative when P_s is positive, where P_u is the power flowing into the nonlinear reactance at ω_u . This means that the external circuit receives a net power at ω_u when the signal power is flowing into the nonlinear reactance. If $\omega_u \gg \omega_s$, the output power at ω_u can be made large compared to the available input power at ω_s , thus providing power amplification. Because of the frequency shift, amplifiers of this type are called up-converters. The discussion in Sections 7.1, 7.2 and 7.4 does not apply to up-converters since input and output frequencies are different in the latter. Combining (7.91) and (7.92), we have

$$P_p + P_s = -P_u \{(\omega_p/\omega_u) + (\omega_s/\omega_u)\} = -P_u$$

where $\omega_u = \omega_p + \omega_s$ is used. The above equation shows that the output

power is made up of the signal and pump powers as expected from the lossless assumption.

Next, suppose that the nonlinear reactance is reactively terminated at all frequencies except ω_p , ω_s , and $\omega_l = \omega_p - \omega_s$, then the Manley-Rowe relations become

$$(P_s/\omega_s) - (P_l/\omega_l) = 0 \quad (7.93)$$

$$(P_p/\omega_p) + (P_l/\omega_l) = 0 \quad (7.94)$$

where $P_{-1,1}$ is indicated by P_l . If the external circuit is passive at ω_l , P_l must be negative, and hence P_s becomes negative also from (7.93). In other words, a net power is flowing out from the nonlinear reactance at ω_s , which is given by

$$-P_s = -P_l(\omega_s/\omega_l)$$

If we adjust the external circuit properly so that $|P_l|$ becomes large, then $|P_s|$ will also be large. Therefore, as far as the circuit at ω_s is concerned, the nonlinear reactance is behaving as a negative resistance, and an amplifier can be built utilizing this property. Such an amplifier is called a negative resistance parametric amplifier. For negative resistance parametric amplifiers, the output and input frequencies are identical, and all the discussions given in the previous sections are directly applicable. The circuit absorbing power at ω_l , which is an essential part of the amplification process, is called the idler circuit, and ω_l the idler frequency. This name is appropriate since the power at ω_l is not directly utilized. Combining (7.93) and (7.94), we have

$$-(P_s + P_l) = -\{(\omega_s/\omega_l) + 1\} P_l = \{(\omega_s/\omega_l) + 1\} P_p(\omega_l/\omega_p) = P_p$$

which indicates that the total power flowing out from the nonlinear reactance at ω_s and ω_p is supplied by the pump, as expected.

7.7 Negative Resistance Parametric Amplifiers

In this section, we shall study negative resistance parametric amplifiers which utilize the variable junction capacitance of a diode. Let v be the voltage across the junction capacitance and q be the stored charge. The relation between v and q is given by (7.79). Suppose a large voltage v_p having a fundamental frequency ω_p is applied to the junction and the corresponding charge is given by q_p . Then

$$v_p = f(q_p)$$

The charge q_p is a periodic function of time having a fundamental frequency ω_p . Suppose a small signal voltage at ω_s is superimposed which changes the junction voltage v_p and charge q_p to $v_p + \delta v$ and $q_p + \delta q$, respectively. Then

$$v_p + \delta v = f(q_p + \delta q) \quad (7.95)$$

from which we obtain

$$\delta v = \frac{\partial f(q_p)}{\partial q_p} \delta q \quad (7.96)$$

Since q_p is a periodic function of time and f is a single-valued function, $\partial f/\partial q_p$ is also a periodic function of time having the same fundamental frequency ω_p . Choosing the time origin properly, we can write $\partial f/\partial q_p$ in the form

$$\frac{\partial f(q_p)}{\partial q_p} = \frac{1}{K_0} + \frac{1}{K_1} \cos \omega_p t + \frac{1}{K_2} \cos 2\omega_p t + \dots \quad (7.97)$$

which is simply the Fourier series expansion of the periodic function. Note that $K_0, K_1, K_2, \dots$ have the dimension of capacitance.

The functions δv and δq can be expressed as follows:

$$\delta v = \sum' (v_{m,n}/\sqrt{2}) e^{j(m\omega_s + n\omega_p)t} \quad (7.98)$$

$$\delta q = \sum' (q_{m,n}/\sqrt{2}) e^{j(m\omega_s + n\omega_p)t} \quad (7.99)$$

where

$$v_{m,n} = v_{-m,-n}^*, \quad q_{m,n} = q_{-m,-n}^* \quad (7.100)$$

and $\sum'$ indicates the summation with respect to m and n from $-\infty$ to $+\infty$, excluding those terms for which $m=0$. The junction current δi corresponding to δv is given by

$$\delta i = \sum' (i_{m,n}/\sqrt{2}) e^{j(m\omega_s + n\omega_p)t} \quad (7.101)$$

where

$$i_{m,n} = j(m\omega_s + n\omega_p) q_{m,n} \quad (7.102)$$

Suppose that the external impedance is properly adjusted so that the current can be neglected at all frequencies expressible in the form $|m\omega_s + n\omega_p|$ except ω_s and $\omega_l = \omega_p - \omega_s$. Then only four of the $q_{m,n}$'s ($m \neq 0$) remain nonvanishing: $q_{1,0}$, $q_{-1,0}$, $q_{-1,1}$, and $q_{1,-1}$. Because of (7.100), only two of them are independent so let us try to write the voltages at ω_s and ω_l in terms of these independent variables. If we substitute (7.97), (7.98), and (7.99) into (7.96) and compare the terms with $\exp\{j\omega_s t\}$ on both sides, we

obtain

$$v_{1,0} = (1/K_0) q_{1,0} + (1/2K_1) q_{1,-1}$$

Similarly, comparing the terms with $\exp\{j(\omega_s - \omega_p)t\}$, we have

$$v_{1,-1} = (1/K_0) q_{1,-1} + (1/2K_1) q_{1,0}$$

Using (7.102), the above result can be rewritten as follows:

$$v_s = \frac{1}{j\omega_s K_0} i_s - \frac{1}{j\omega_l (2K_1)} i_l^* \quad (7.103)$$

$$v_l^* = -\frac{1}{j\omega_l K_0} i_l^* + \frac{1}{j\omega_s (2K_1)} i_s \quad (7.104)$$

where $v_{1,0}$ and $i_{1,0}$ are indicated by v_s and i_s since they are the voltage and current at ω_s . Similarly, $v_{1,-1}$ and $i_{-1,1}$ are indicated by v_l and i_l , respectively. In addition, there may be nonzero voltages at other frequencies. However, since the corresponding currents are zero by hypothesis, no net power transfer takes place, and hence we do not bother with them.

From (7.103) and (7.104), we see that an equivalent circuit of the junction is given by Fig. 7.25, in which the ω_s and ω_p circuits are drawn separately

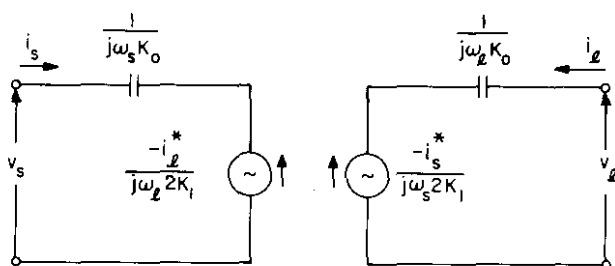


Fig. 7.25. Equivalent circuit of a parametric junction.

and $1/j\omega_s K_0$ and $1/j\omega_l K_0$ are the impedances due to the series capacitance K_0 which play no essential role in the parametric amplification. The terms containing $2K_1$ which are drawn in the form of voltage sources are responsible for the amplification. Figure 7.25 represents the diode junction only, but in practice the junction current flows through the semiconductor part which acts as a small, but finite, series resistance r . To take its effect into account let us define $\tilde{Q}_s$ and $\tilde{Q}_l$ by

$$\tilde{Q}_s = (1/\omega_s 2K_1 r), \quad \tilde{Q}_l = (1/\omega_l 2K_1 r) \quad (7.105)$$

respectively. Remember that the Q of an ordinary capacitance C is defined by $1/\omega Cr$, and that the Q 's defined above have $2K_1$ in place of C , where $2K_1$ represents a capacitance component changing with time, it is therefore logical to call $\tilde{Q}_s$ and $\tilde{Q}_l$ the dynamic Q of the variable capacitance at ω_s and ω_l , respectively. Suppose that a load impedance is connected to the ω_l circuit and let Z_l represent the total impedance due to $1/j\omega_l K_0$ and the load impedance in series, then the ω_l circuit becomes the series connection of r , Z_l , a voltage source

$$-\frac{i_s^*}{j\omega_s 2K_1} = j\tilde{Q}_s r i_s^*$$

and a possible noise voltage v_n as shown in Fig. 7.26. If we assume that both

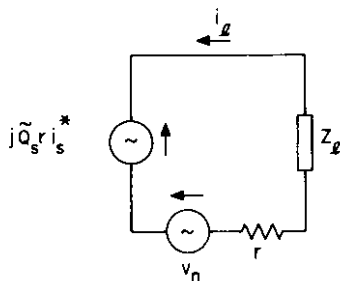


Fig. 7.26. Equivalent idler circuit.

r and $R_l = \text{Re}\{Z_l\}$ are at temperature T and that v_n represents the thermal noise from these resistances, $\langle v_n^2 \rangle$ is given by

$$\langle v_n^2 \rangle = 4kTB(r + R_l) \quad (7.106)$$

From Fig. 7.26, i_l is calculated to be

$$i_l = -(r + Z_l)^{-1} v_n - j\tilde{Q}_s r (r + Z_l)^{-1} i_s^*$$

The equivalent voltage source in the ω_s circuit shown in Fig. 7.25 is, therefore, given by

$$-\frac{i_l^*}{j\omega_l 2K_1} = j\tilde{Q}_l r i_l^* = -j\tilde{Q}_l r (r + Z_l^*)^{-1} v_n^* - \tilde{Q}_l \tilde{Q}_s r^2 (r + Z_l^*)^{-1} i_s \quad (7.107)$$

The second term on the right-hand side of (7.107) expresses the voltage developed across the impedance $-\tilde{Q}_l \tilde{Q}_s r^2 (r + Z_l^*)^{-1}$ when i_s flows. Con-

sequently, the equivalent circuit looking into the diode at ω_s becomes as shown in Fig. 7.27, where $\langle e_n^2 \rangle$ is given by

$$\langle e_n^2 \rangle = 4kTBr \quad (7.108)$$

In order to obtain amplification, the real part of the impedance must be negative:

$$\text{Re} \{ r - \tilde{Q}_i \tilde{Q}_s r^2 (r + Z_i^*)^{-1} \} < 0 \quad (7.109)$$

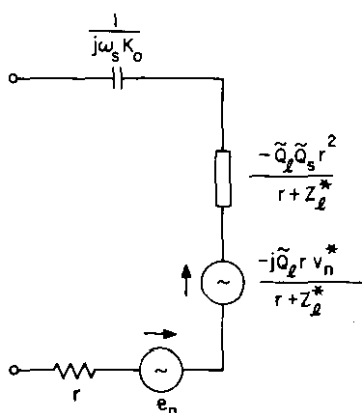


Fig. 7.27. Equivalent circuit of the diode at ω_s .

The second term in the bracket becomes maximum when Z_i^* is equal to zero. Setting $Z_i^* = 0$, (7.109) reduces to

$$r(1 - \tilde{Q}_i \tilde{Q}_s) < 0 \quad (7.110)$$

If ω_i and ω_s become large and $\tilde{Q}_i \tilde{Q}_s$ becomes less than unity, (7.110), and hence (7.109), cannot be satisfied and no amplification is expected.

Let us next calculate M_{opt} . Since e_n and v_n are independent, the mean square value of the open-circuit noise voltage V_n is given by

$$\langle V_n^2 \rangle = \langle e_n^2 \rangle + \langle v_n^2 \rangle |j\tilde{Q}_i r|^2 |r + Z_i|^2$$

Substituting (7.107) and (7.108) and writing Z_i in the form $R_i + jX_i$, we have

$$\langle V_n^2 \rangle = 4kTBr + 4kTB(r + R_i) \tilde{Q}_i^2 r^2 \{ (r + R_i)^2 + X_i^2 \}^{-1} \quad (7.111)$$

The exchangeable noise power P_e is given by $\langle V_n^2 \rangle$ divided by four times the real part of the equivalent impedance of the diode shown in Fig. 7.27.

Since the real part is given by

$$r - \tilde{Q}_l \tilde{Q}_s r^2 (r + R_l) \{(r + R_l)^2 + X_l^2\}^{-1}$$

we have

$$M_{\text{opt}} = \frac{-P_e}{kT_l B} = \frac{T}{T_l} \frac{\tilde{Q}_l^2 r (r + R_l) + (r + R_l)^2 + X_l^2}{\tilde{Q}_l \tilde{Q}_s r (r + R_l) - \{(r + R_l)^2 + X_l^2\}} \quad (7.112)$$

This is the M_{opt} for fixed R_l and X_l . If X_l is allowed to change, the M_{opt} reaches a minimum when $X_l^2 = 0$, since the numerator increases and the denominator decreases as X_l^2 increases from zero. This minimum value of M_{opt} is given by

$$M_{\text{opt}} = \frac{T}{T_l} \frac{\tilde{Q}_l^2 r + (r + R_l)}{\tilde{Q}_l \tilde{Q}_s r - (r + R_l)}$$

Furthermore, if R_l is allowed to change, the best value of M_{opt} is attained when $R_l = 0$ since the numerator increases and the denominator decreases as R_l increases from zero. The best M_{opt} is given by

$$M_{\text{opt}} = \frac{T}{T_l} \frac{\tilde{Q}_l^2 + 1}{\tilde{Q}_s \tilde{Q}_l - 1} \quad (7.113)$$

which is obtainable when, and only when,

$$R_l = 0, \quad X_l = 0 \quad (7.114)$$

or equivalently, when the external idler circuit is an inductance which just cancels $1/j\omega_l K_0$. As in the case of tunnel diode amplifiers, M_{opt} can be realized by connecting the generator and load to the diode through a circulator. The gain can be adjusted to a desired value, by inserting a proper lossless circuit between the circulator and the diode.

If $T = T_l$, Eq. (7.113) is a function of $\tilde{Q}_s$ and $\tilde{Q}_l$ only. Since $\tilde{Q}_l$ is determined from $\tilde{Q}_s$ and ω_l/ω_s , $M_{\text{opt}} \times T_l$ can be plotted as a function of $\tilde{Q}_s$ with ω_l/ω_s as a parameter. Figure 7.28 shows such a plot. From this, we see that, for fixed $\tilde{Q}_s$, there is an optimum value of ω_l/ω_s which gives the minimum M_{opt} .

Note that T is the diode temperature, and that a further improvement of M_{opt} is expected if the diode is cooled. For example, the value of $M_{\text{opt}} \times T_l$ becomes 77.2/290 times the value shown in Fig. 7.28 when the diode is cooled to the temperature of liquid nitrogen, and 4.2/290 times when cooled to the temperature of liquid helium. It should be pointed out that $\tilde{Q}_s$ of a diode is a function of temperature, and hence the room temperature value must not be used in the refrigerated cases. Furthermore, it should be

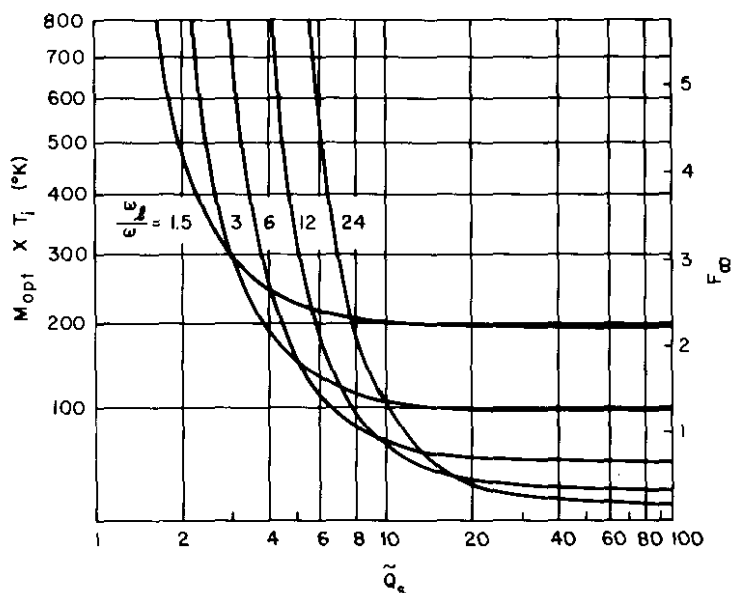


Fig. 7.28. $M_{opt} \times T_i$ versus $\tilde{Q}_s$ with ω_l/ω_s as a parameter.

pointed out that noise originating in the circulator may predominate in such a low noise application unless the circulator is also cooled.

PROBLEMS

- 7.1 Obtain the canonical form of a triode assuming the equivalent circuit as shown in Fig. 7.29.
- 7.2 Obtain the condition under which the triode can provide a positive resistance with the noise temperature below T .

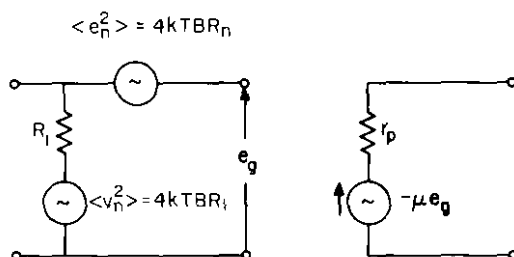


Fig. 7.29. Equivalent circuit of a triode.

- 7.3 Calculate the available and transducer gains of the amplifier shown in Fig. 7.30, assuming $R_g = R_L = 50 \Omega$ and $R = 52 \Omega$. Note: A tunnel diode with negligible series resistance r can be represented by the circuit inside the dotted line.
- 7.4 Suppose that an amplifier A has 50Ω input and output impedances, and let G be the transducer gain when the generator and load impedances are $(50 + j0) \Omega$. Inserting the amplifier inside the dotted line in Fig. 7.30 in front of amplifier A , calculate the overall available and transducer gains of the cascaded amplifier. Also, calculate the available and transducer gains of each component amplifier in the cascade connection.

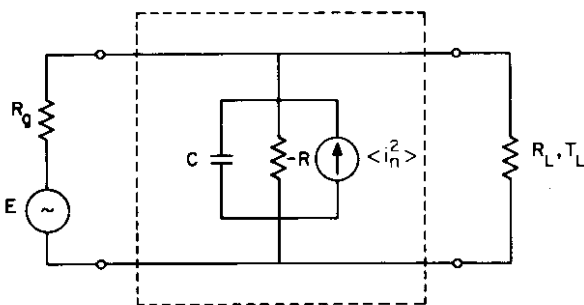


Fig. 7.30. An amplifier for which $M_e = M_{\text{opt}}$ but M is poor.

- 7.5 The exchangeable gain G_e of an amplifier is defined by the ratio of the exchangeable signal power at the output port to the available signal power from the generator, provided that the generator impedance has a positive real part. The exchangeable noise figure F_e is the ratio of the exchangeable noise power at the output port to kT_iBG_e , when the noise temperature T_i of the generator is specified to be 290°K . The exchangeable noise measure M_e is defined by

$$M_e = \frac{F_e - 1}{1 - (1/G_e)}$$

Calculate the possible range of M_e following the discussion for M given in Section 7.2, and show that the smallest positive value of M_e is equal to M_{opt} . Also show that $M \geq M_e$ whenever M is positive. This means that the evaluation of amplifier noise performance by M is generally more critical than that obtained from M_e .

- 7.6 Calculate M_e and M for the amplifier shown in Fig. 7.30, and show that M_e of this amplifier is always equal to M_{opt} , regardless of the values of R_g , R_L , R , and C . Although this type of amplifier can give a large available gain over a wide bandwidth and the noise performance expressed by M_e is optimum, it is generally not considered to be a good practical amplifier because of its low transducer gain and poor M .
- 7.7 Calculate the operating noise temperature of a parametric amplifier assuming the output power is taken at the idler frequency.
- 7.8 Prove that $\tilde{Q}$ is invariant to nonsingular, lossless transformations.

- 7.9 Let γ_1 and γ_2 be the eigenvalues of $Z + Z^+$ of a two-port network, and prove that

$$\Gamma = (\gamma_1/|\gamma_1|) + (\gamma_2/|\gamma_2|) \quad (7.115)$$

is invariant to nonsingular, lossless, reciprocal transformations.

- 7.10 Prove that the two circuits shown in Fig. 7.31 have the same unilateral gain U but that no lossless reciprocal network can transform one to the other. *Note:* For these

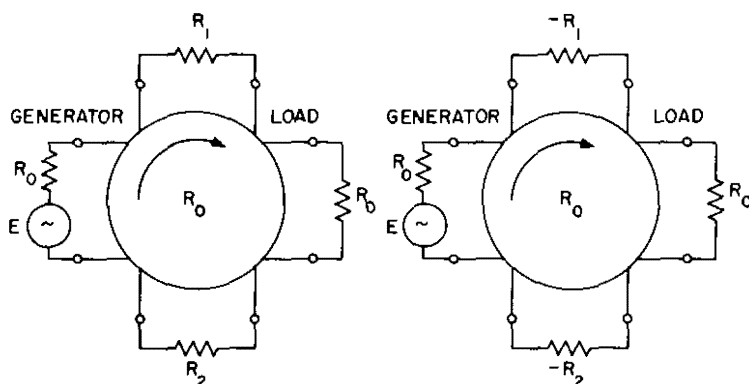


Fig. 7.31. Two circuits which have the same unilateral gain but cannot be transformed from one to the other by any lossless reciprocal network.

two circuits, the values of (7.115) are different. Consequently, the unilateral gain U is not a unique invariant to nonsingular lossless reciprocal transformations.

- 7.11 Prove that the condition under which both input and output ports of a two-port network ($|S_{12}S_{21}| \neq 0$) can be matched simultaneously keeping the real parts of the source and load impedances positive is given by

$$|S_{11}|^2 + |S_{22}|^2 + 2|S_{12}S_{21}| < 1 + |S_{12}S_{21} - S_{11}S_{22}|^2$$

When $|S_{12}S_{21}| = 0$, the same condition is given by

$$|S_{11}| < 1, \quad |S_{22}| < 1$$



CHAPTER 8

ELECTRON BEAMS

A stream of electrons in free space is called an electron beam. It plays an essential role in most microwave electron tubes such as klystrons and traveling wave tubes. In this chapter we shall develop a small signal theory for a simplified electron beam model in order to gain physical insight into the behavior of the actual beam. It must be pointed out, however, that the small signal assumption imposes a rather stringent restriction on the applicability of the theory. For example, we cannot use the present discussion to estimate some of the important parameters of microwave tubes, such as efficiency and maximum output power.

We shall first construct the small signal model of electron beams and study the properties of possible waves propagating along the model. Next, the interaction between the electron beam and electrode gaps, as well as the continuous interaction with a slow wave structure, will be discussed in detail together with an explanation of the signal amplification process in microwave tubes. The last section is devoted to a discussion of amplifier noise performance that can be expected with a given electron beam.

8.1 An Electron Beam Model

To construct a simple yet useful electron beam model let us assume that the electrons in the beam move in the z direction only. The transverse motion can be restricted, for instance, by applying a high static magnetic field in the z direction.

Corresponding to the electron motion in the z direction, there is current

$i_z(\mathbf{r}, t)$, where $\mathbf{r}$ represents position (x, y, z) . Maxwell's equations then become

$$\nabla \times \mathbf{E}(\mathbf{r}, t) = -\mu \frac{\partial}{\partial t} \mathbf{H}(\mathbf{r}, t) \quad (8.1)$$

$$\nabla \times \mathbf{H}(\mathbf{r}, t) = \mathbf{k}_z i_z(\mathbf{r}, t) + \epsilon \frac{\partial}{\partial t} \mathbf{E}(\mathbf{r}, t) \quad (8.2)$$

Electrons in the beam are accelerated by the electric field $E_z(\mathbf{r}, t)$. The equation of motion is given by

$$\frac{d}{dt} v(\mathbf{r}, t) = -\frac{e}{m} E_z(\mathbf{r}, t) \quad (8.3)$$

where m is the electron mass 9.108×10^{-31} kg, e is the magnitude of electron charge 1.602×10^{-19} coulomb, $v(\mathbf{r}, t)$ is the electron velocity, and d/dt indicates the total derivative with respect to time. Let $\partial/\partial t$ indicate the partial derivative with respect to time; i.e., the derivative taken at a fixed point $\mathbf{r}$ as usual, then we have

$$\frac{d}{dt} v(\mathbf{r}, t) = \frac{\partial}{\partial t} v(\mathbf{r}, t) + \frac{dz}{dt} \frac{\partial}{\partial z} v(\mathbf{r}, t)$$

Substituting the above equation into (8.3) gives

$$\frac{\partial}{\partial t} v(\mathbf{r}, t) + v(\mathbf{r}, t) \frac{\partial}{\partial z} v(\mathbf{r}, t) = -\frac{e}{m} E_z(\mathbf{r}, t) \quad (8.4)$$

where $dz/dt = v(\mathbf{r}, t)$ is used.

The equation of continuity, expressing the conservation of charge, is given by

$$\frac{\partial}{\partial z} i_z(\mathbf{r}, t) = -\frac{\partial}{\partial t} \rho(\mathbf{r}, t) \quad (8.5)$$

The current density is equal to the product of charge density and electron velocity, i.e.,

$$i_z(\mathbf{r}, t) = v(\mathbf{r}, t) \rho(\mathbf{r}, t) \quad (8.6)$$

All the quantities introduced above may have two distinct components, the stationary and alternating components. Since the alternating components are sufficiently small under the small signal assumption, the product of two alternating components can be neglected compared to the product of corresponding stationary and alternating components. A consistent model

can be obtained by taking into account the stationary and fundamental frequency components only while neglecting all the possible higher harmonics. This reasoning enables us to assume the following relations:

$$\begin{aligned}\mathbf{E}(\mathbf{r}, t) &= \mathbf{E}_0(\mathbf{r}) + \operatorname{Re} \{ \sqrt{2} \mathbf{E}(\mathbf{r}) e^{j\omega t} \} \\ \mathbf{H}(\mathbf{r}, t) &= \mathbf{H}_0(\mathbf{r}) + \operatorname{Re} \{ \sqrt{2} \mathbf{H}(\mathbf{r}) e^{j\omega t} \} \\ i_z(\mathbf{r}, t) &= i_{z0}(\mathbf{r}) + \operatorname{Re} \{ \sqrt{2} i_z(\mathbf{r}) e^{j\omega t} \} \\ v(\mathbf{r}, t) &= v_0(\mathbf{r}) + \operatorname{Re} \{ \sqrt{2} v(\mathbf{r}) e^{j\omega t} \} \\ \varrho(\mathbf{r}, t) &= \varrho_0(\mathbf{r}) + \operatorname{Re} \{ \sqrt{2} \varrho(\mathbf{r}) e^{j\omega t} \}\end{aligned}\quad (8.7)$$

Substituting (8.7) into (8.1) and (8.2), we obtain

$$\nabla \times \mathbf{E}_0(\mathbf{r}) = 0 \quad (8.8)$$

$$\nabla \times \mathbf{H}_0(\mathbf{r}) = \mathbf{k} i_{z0}(\mathbf{r}) \quad (8.9)$$

$$\nabla \times \mathbf{E}(\mathbf{r}) = -j\omega\mu\mathbf{H}(\mathbf{r}) \quad (8.10)$$

$$\nabla \times \mathbf{H}(\mathbf{r}) = \mathbf{k} i_z(\mathbf{r}) + j\omega\varepsilon\mathbf{E}(\mathbf{r}) \quad (8.11)$$

Similarly from (8.4), we have

$$v_0(\mathbf{r}) \frac{\partial v_0(\mathbf{r})}{\partial z} = -\frac{e}{m} E_{z0}(\mathbf{r}) \quad (8.12)$$

$$j\omega v(\mathbf{r}) + v_0(\mathbf{r}) \frac{\partial v(\mathbf{r})}{\partial z} + v(\mathbf{r}) \frac{\partial v_0(\mathbf{r})}{\partial z} = -\frac{e}{m} E_z(\mathbf{r}) \quad (8.13)$$

where the product term of two alternating components, $v(\mathbf{r}) \{ \partial v(\mathbf{r}) / \partial z \}$, has been neglected. We obtain from (8.5)

$$\frac{\partial i_{z0}(\mathbf{r})}{\partial z} = 0 \quad (8.14)$$

$$\frac{\partial i_z(\mathbf{r})}{\partial z} = -j\omega\varrho(\mathbf{r}) \quad (8.15)$$

and from (8.6)

$$i_{z0}(\mathbf{r}) = v_0(\mathbf{r}) \varrho_0(\mathbf{r}) \quad (8.16)$$

$$i_z(\mathbf{r}) = v_0(\mathbf{r}) \varrho(\mathbf{r}) + v(\mathbf{r}) \varrho_0(\mathbf{r}) \quad (8.17)$$

in which another product term of two alternating components, $v(\mathbf{r}) \varrho(\mathbf{r})$, has been neglected. Equations (8.8)–(8.17) specify the behavior of our electron beam model.

Let us now study the transmission power along the model. To do so, we calculate $\nabla \cdot \mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r})$ following the discussion of Section 2.2. From

(8.10) and (8.11), we have

$$\nabla \cdot \mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r}) = -j\omega\mu\mathbf{H}(\mathbf{r}) \cdot \mathbf{H}^*(\mathbf{r}) + j\omega\varepsilon\mathbf{E}(\mathbf{r}) \cdot \mathbf{E}^*(\mathbf{r}) - E_z(\mathbf{r}) i_z^*(\mathbf{r}) \quad (8.18)$$

Using (8.13) and (8.17), the last term on the right-hand side can be written in the form

$$E_z(\mathbf{r}) i_z^*(\mathbf{r}) = -\frac{m}{e} \left\{ j\omega\varrho_0(\mathbf{r}) v(\mathbf{r}) v^*(\mathbf{r}) + j\omega v(\mathbf{r}) v_0(\mathbf{r}) \varrho^*(\mathbf{r}) + i_z^*(\mathbf{r}) \frac{\partial}{\partial z} v_0(\mathbf{r}) v(\mathbf{r}) \right\} \quad (8.19)$$

The second term in the bracket is equal to

$$v_0(\mathbf{r}) v(\mathbf{r}) \frac{\partial i_z^*(\mathbf{r})}{\partial z}$$

as is easily seen from (8.15). This can be combined with the last term in the bracket of (8.19) to give

$$\frac{\partial}{\partial z} v_0(\mathbf{r}) v(\mathbf{r}) i_z^*(\mathbf{r})$$

Substituting (8.19) and this result into (8.18), and remembering that $k \cdot \nabla$ is the partial derivative with respect to z , we obtain

$$\begin{aligned} \nabla \cdot \{ \mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r}) - (m/e) \mathbf{k} v_0(\mathbf{r}) v(\mathbf{r}) i_z^*(\mathbf{r}) \} \\ = j\omega \{ \varepsilon \mathbf{E}(\mathbf{r}) \cdot \mathbf{E}^*(\mathbf{r}) - \mu \mathbf{H}(\mathbf{r}) \cdot \mathbf{H}^*(\mathbf{r}) + (m/e) \varrho_0(\mathbf{r}) v(\mathbf{r}) v^*(\mathbf{r}) \} \end{aligned}$$

Let us integrate this equation over a volume V containing the electron beam as shown in Fig. 8.1. The left-hand side can be converted to a surface integral over S resulting in

$$\begin{aligned} \int \{ \mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r}) - (m/e) \mathbf{k} v_0(\mathbf{r}) v(\mathbf{r}) i_z^*(\mathbf{r}) \} \cdot \mathbf{n} dS \\ = j\omega \int \{ \varepsilon \mathbf{E}(\mathbf{r}) \cdot \mathbf{E}^*(\mathbf{r}) - \mu \mathbf{H}(\mathbf{r}) \cdot \mathbf{H}^*(\mathbf{r}) + (m/e) \varrho_0(\mathbf{r}) v(\mathbf{r}) v^*(\mathbf{r}) \} dv \end{aligned} \quad (8.20)$$

Since the right-hand side is pure imaginary, the real part of (8.20) gives

$$\text{Re} \int \mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r}) \cdot \mathbf{n} dS + \text{Re} \int (-1)(m/e) \mathbf{k} v_0(\mathbf{r}) v(\mathbf{r}) i_z^*(\mathbf{r}) \cdot \mathbf{n} dS = 0 \quad (8.21)$$

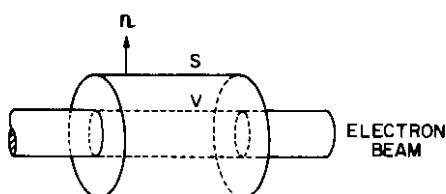


Fig. 8.1. Region of integration.

The first term on the left-hand side is the electromagnetic power P flowing out from the volume V through the surface S ;

$$P = \text{Re} \int \mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r}) \cdot \mathbf{n} dS \quad (8.22)$$

The second term on the left-hand side of (8.21) is related to the motion of electrons. Since (8.21) shows that the sum of these two terms is always equal to zero, let us interpret the second term as the kinetic power P_k carried out by the electrons through S . With this interpretation, (8.21) indicates the conservation of energy: The electromagnetic power flowing out from a volume V is equal to the kinetic power flowing into the same volume. Power is expressed ordinarily as the product of voltage and current. If we define the kinetic voltage of the electron beam at $\mathbf{r}$ by

$$V_k(\mathbf{r}) = - (m/e) v_0(\mathbf{r}) v(\mathbf{r}) \quad (8.23)$$

then the kinetic power is given by

$$P_k = \text{Re} \int V_k(\mathbf{r}) i_z^*(\mathbf{r}) \mathbf{k} \cdot \mathbf{n} dS \quad (8.24)$$

Just as we interpret $\mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r})$ as the electromagnetic power density, the kinetic voltage times electron current density $V_k(\mathbf{r}) i_z^*(\mathbf{r})$ can be considered as the kinetic power density of the electron beam at $\mathbf{r}$.

The first two terms on the right-hand side of (8.20),

$$\int \epsilon \mathbf{E}(\mathbf{r}) \cdot \mathbf{E}^*(\mathbf{r}) dv, \quad \int \mu \mathbf{H}(\mathbf{r}) \cdot \mathbf{H}^*(\mathbf{r}) dv$$

represent twice the average stored electric and magnetic energies, respectively. The third term

$$\int (m/e) \varrho_0(\mathbf{r}) v(\mathbf{r}) v^*(\mathbf{r}) dv$$

can be interpreted as twice the average stored kinetic energy. Since e is positive and $\varrho_0(\mathbf{r})$ is negative, expressing the electron charge density, this term has the same sign as the stored magnetic energy in (8.20).

8.2 Space Charge Waves

Following the discussion of waveguides in Chapter 3, let us assume that the microwave components of various quantities vary with distance exponentially; i.e.,

$$\begin{aligned} \mathbf{E}(\mathbf{r}) &= (\mathbf{E}_t + \mathbf{k}E_z) e^{-\gamma z}, & \mathbf{H}(\mathbf{r}) &= (\mathbf{H}_t + \mathbf{k}H_z) e^{-\gamma z} \\ i_z(\mathbf{r}) &= i_z e^{-\gamma z}, & v(\mathbf{r}) &= v e^{-\gamma z}, & \varrho(\mathbf{r}) &= \varrho e^{-\gamma z} \end{aligned} \quad (8.25)$$

where $\mathbf{E}_t$, E_z , $\mathbf{H}_t$, H_z , i_z , v , and ϱ are independent of z but may be functions of transverse position. We further assume that the stationary component of electron velocity has the same value everywhere over the cross section of the electron beam. This assumption is approximately satisfied after the electrons have been accelerated by a high static voltage. Furthermore, let us assume that the stationary component of the electron velocity does not depend on z ; i.e., no static electric field is applied in the region of interest. Then, (8.13) becomes

$$j\omega v - \gamma v_0 v = - (e/m) E_z \quad (8.26)$$

Similarly, from (8.15) and (8.17) we have

$$- \gamma i_z = - j\omega \varrho \quad (8.27)$$

$$i_z = v_0 \varrho + v \varrho_0 \quad (8.28)$$

From (8.26), v is given by

$$v = - (e/m) (j\omega - \gamma v_0)^{-1} E_z \quad (8.29)$$

By solving (8.27) and (8.28) ϱ can be expressed in terms of v . Substituting this result into (8.29), we obtain

$$\varrho = \gamma \varrho_0 (j\omega - \gamma v_0)^{-1} v = - \gamma \varrho_0 (e/m) (j\omega - \gamma v_0)^{-2} E_z \quad (8.30)$$

Finally, by substituting (8.30) into (8.27), we obtain i_z in terms of E_z :

$$i_z = - j\omega \varrho_0 (e/m) (j\omega - \gamma v_0)^{-2} E_z \quad (8.31)$$

If E_z is equal to zero, i_z is also zero from (8.31), and Maxwell's equations (8.10) and (8.11) assume the same forms as in the case in which there is no electron beam. This means that TE waves propagate independently of the

electron beam. To study the effect of the electron beam, however, we concentrate on the case in which E_z is not equal to zero.

Beginning with (8.10) and (8.11) in place of (3.52) and (3.53), and proceeding as we did in Section 3.4, we find that (3.58), (3.59), (3.60), and (3.62) are also valid for the present discussion without modification. On the other hand, (3.57) and (3.61) are replaced by

$$\nabla \times \mathbf{H}_t = j\omega\epsilon \{1 + \omega_p^2 (j\omega - \gamma v_0)^{-2}\} \mathbf{k} E_z \quad (8.32)$$

and

$$\nabla \cdot \mathbf{E}_t = \{1 + \omega_p^2 (j\omega - \gamma v_0)^{-1}\} \gamma E_z \quad (8.33)$$

respectively, where use is made of (8.31) and

$$\omega_p^2 = -(\rho_0/\epsilon)(e/m) \quad (8.34)$$

Taking $\nabla \times$ (3.60) and rearranging the left-hand side, we have

$$\gamma \mathbf{k} \nabla \cdot \mathbf{E}_t + \mathbf{k} \nabla \cdot \nabla E_z = j\omega\mu \nabla \times \mathbf{H}_t$$

Substituting (8.32) and (8.33) into this expression gives

$$\nabla \cdot \nabla E_z + (\omega^2 \epsilon \mu + \gamma^2) \{1 + \omega_p^2 (j\omega - \gamma v_0)^{-2}\} E_z = 0 \quad (8.35)$$

A perfect conductor wall gives the boundary condition $E_z = 0$ independent of H_z . In a similar manner, suppose that a boundary condition for E_z is given independent of H_z , then E_z can be obtained which satisfies (8.35) under the boundary condition while assuming $H_z = 0$. Once E_z is obtained, $\mathbf{H}_t$ and $\mathbf{E}_t$ are given by

$$\mathbf{H}_t = -j\omega\epsilon (\omega^2 \epsilon \mu + \gamma^2)^{-1} \mathbf{k} \times \nabla E_z \quad (8.36)$$

$$\mathbf{E}_t = -\gamma (\omega^2 \epsilon \mu + \gamma^2)^{-1} \nabla E_z \quad (8.37)$$

where (3.58) and (3.60) have been used. Furthermore, v , ρ , and i_z can be calculated from (8.29), (8.30), and (8.31), respectively. The problem is thus reduced to the solution of (8.35) under the appropriate boundary condition. It is difficult to give a general discussion of the eigenvalue problem represented by (8.35). Hence, we shall discuss a few particularly simple examples and imply what kinds of solutions may exist in more general cases.

First, let us consider a one-dimensional model. There are no field variations in the x and y directions, and the first term on the left-hand side of (8.35) becomes zero thereby making the second term vanish. If we assume that v_0 is equal to zero, the nontrivial solution E_z can exist when $\omega = \omega_p$. This solution represents a free-running oscillation due to the repulsion force between electronic charge and the inertia of electrons. Since free electrons

in a plasma can oscillate in this way, ω_p is called the plasma angular frequency.

For the electron beam, v_0 is not equal to zero. In order to obtain non-trivial solutions from (8.35), either

$$\omega^2 \epsilon \mu + \gamma^2 = 0 \quad (8.38)$$

or

$$\{1 + \omega_p^2 (j\omega - \gamma v_0)^{-2}\} = 0 \quad (8.39)$$

has to be satisfied. From (8.32) and the one-dimensional assumption, E_z is equal to zero unless (8.39) is satisfied. Therefore, we have only to investigate (8.39). There are two different values of γ which satisfy (8.39):

$$\gamma_+ = j(\omega - \omega_p)/v_0 \quad (8.40)$$

$$\gamma_- = j(\omega + \omega_p)/v_0 \quad (8.41)$$

This means that there are two different waves, one slightly faster and the other slightly slower than v_0 , corresponding to γ_+ and γ_- , respectively; they are called the fast and slow waves.

Since $\nabla E_z = 0$ for the one-dimensional model, both $\mathbf{E}_t$ and $\mathbf{H}_t$ are equal to zero, and no electromagnetic power is transmitted. Substituting (8.40) and (8.41) into (8.31), we have

$$i_z = -j\omega \epsilon E_z \quad (8.42)$$

for both fast and slow waves. This indicates that the electronic and displacement currents cancel each other and explains why $\mathbf{H}_t$ becomes zero.

When (8.38) is satisfied, even if E_z is equal to zero, $\mathbf{E}_t$ and $\mathbf{H}_t$ may have nonzero magnitudes, corresponding to the plane wave in free space.

Let us next consider an electron beam moving in the axial direction, which fills the inside of a waveguide. We assume that ω_p^2 is constant over the waveguide cross section and solve (8.35) under the boundary condition $E_z = 0$. From the discussion of waveguides, the eigenvalue problem

$$\nabla^2 E_z + k^2 E_z = 0 \quad (\text{in } S)$$

$$E_z = 0 \quad (\text{on } L)$$

has an infinite number of solutions, where S indicates the waveguide cross section and L the boundary. Let E_{zn} be the n th solution and k_n^2 the corresponding eigenvalue. Then E_{zn} itself serves as a solution for (8.35) and the corresponding γ is obtained from

$$k_n^2 = (\omega^2 \epsilon \mu + \gamma^2) \{1 + \omega_p^2 (j\omega - \gamma v_0)^{-2}\} \quad (8.43)$$

This is a fourth-order algebraic equation which gives four values of γ . Of these four values, we are interested in those which have a strong interaction with the electron beam. Since their velocities must be approximately equal to that of the electrons, we assume

$$\gamma_{\pm} \simeq j(\omega \mp \omega_q)/v_0 \quad (8.44)$$

where $\omega \gg \omega_q$. Substituting into (8.43), we then have

$$k_n^2 \simeq -(\omega/v_0)^2 \{1 - (\omega_p/\omega_q)^2\}$$

where $\omega^2 \epsilon \mu$ is neglected compared to $\gamma^2 \simeq -(\omega/v_0)^2$ since $\omega^2 \epsilon \mu$ is equal to ω^2 divided by the square of the velocity of light, and the electron velocity v_0 is far smaller than the velocity of light in ordinary cases. From the above approximate equation, ω_q^2 is given by

$$\omega_q^2 \simeq \omega_p^2 \{\beta_e^2 / (\beta_e^2 + k_n^2)\} \quad (8.45)$$

where

$$\beta_e = (\omega/v_0) \quad (8.46)$$

The phase constant β_e corresponds to the phase velocity v_0 .

Since k_n^2 is positive, from (8.45), we have

$$\omega_p^2 > \omega_q^2 \quad (8.47)$$

A comparison of (8.44) with (8.40) and (8.41) shows that ω_q plays the same role as ω_p did in the one-dimensional model. For these reasons, ω_q is called the reduced plasma angular frequency.

We assumed $\omega \gg \omega_q$ in order to derive (8.45). This assumption does not contradict (8.47) since ϱ_0 is usually small and $\omega \gg \omega_p$.

It is worth noting that since E_z has the same functional form both for the fast and the slow waves, their $\mathbf{E}_t$ and $\mathbf{H}_t$ also have the same functional forms. However, since their γ 's are different, the coefficients in (8.36) and (8.37) have different magnitudes for the two cases. On the other hand, i_z is given by

$$i_z = -j\omega\epsilon(\omega_p^2/\omega_q^2) E_z \quad (8.48)$$

for both cases. Because of (8.47), the electronic current is larger than the displacement current and their difference produces nonzero $\mathbf{H}_t$ in contrast to the one-dimensional model.

Finally, let us consider a more general case in which the electron beam partially fills the waveguide cross section. Since ω_p^2 is now a function of transverse position in the waveguide, it is difficult to solve (8.35) unless some

symmetric configuration is assumed. If the waveguide and the electron beam have concentric circular cross sections and ω_p^2 is constant over the electron beam cross section, then the equation can be solved independently for the regions inside and outside of the electron beam using Bessel functions. These solutions can then be connected to give eigenfunctions which satisfy the continuity of E_z and ∇E_z at the interface. The continuity of ∇E_z is necessary to ensure the continuity of $\mathbf{E}_t$. This analysis (Problem 8.4) shows that there is a pair of propagation constants similar to (8.44) which implies fast and slow waves for each eigenfunction.

Judging from these examples, the fast and slow waves appear to exist in general, and in the remainder of this section we shall assume their existence. They are called space charge waves since the charge of electrons in space is responsible for the wave phenomena.

Let us investigate the relationship between the electromagnetic power P and kinetic power P_k of space charge waves. When the electronic and displacement currents cancel each other, as in the one-dimensional model, P must be zero since no magnetic field exists. In this extreme case, we have

$$(P/P_k) = 0 \quad (8.49)$$

In the other extreme case in which the displacement current is negligibly small compared to the electronic current, (8.11) gives

$$\nabla \times \mathbf{H}_t \simeq k \mathbf{i}_z \quad (8.50)$$

Neglecting $\omega^2 \epsilon \mu$ compared to γ^2 , (8.37) becomes

$$\mathbf{E}_t \simeq -\gamma^{-1} \nabla E_z \quad (8.51)$$

Integrating $\mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r})$ over the waveguide cross section, we have

$$\begin{aligned} \int \mathbf{E}(\mathbf{r}) \times \mathbf{H}^*(\mathbf{r}) \cdot \mathbf{k} dS &= \int \mathbf{E}_t \times \mathbf{H}_t^* \cdot \mathbf{k} dS \simeq -\gamma^{-1} \int (\nabla E_z \times \mathbf{H}_t^*) \cdot \mathbf{k} dS \\ &= -\gamma^{-1} \int \nabla \times E_z \mathbf{H}_t^* \cdot \mathbf{k} dS + \gamma^{-1} \int E_z \nabla \times \mathbf{H}_t^* \cdot \mathbf{k} dS \end{aligned} \quad (8.52)$$

where use is made of (8.51). The first term on the right-hand side of (8.52) becomes a line integral around the inside surface of the waveguide wall by Stokes's theorem and vanishes because $E_z = 0$. The integrand in the second integral on the right-hand side becomes the product of E_z and i_z^* from

(8.50). Thus, the electromagnetic power P is given by

$$P \simeq \text{Re} \left\{ \gamma^{-1} \int E_z i_z^* dS \right\} \quad (8.53)$$

which is independent of z . The kinetic power is calculated from (8.23) and (8.24). In the present case, $v(\mathbf{r})$ in the kinetic voltage is given by (8.29) times $\exp(-\gamma z)$, $i_z^*(\mathbf{r})$ by $i_z^* \exp(\gamma z)$, while $v_0(\mathbf{r})$ is independent of z . Substituting these into (8.24), we have

$$P_k = \text{Re} \left\{ v_0 (j\omega - \gamma v_0)^{-1} \int E_z i_z^* dS \right\} \quad (8.54)$$

which is also independent of z . The integrals in (8.53) and (8.54) are both pure imaginary as is easily seen from (8.31), and the insides of the brackets in both equations become real. The notation Re in (8.53) and (8.54) can therefore be eliminated which makes the ratio between P and P_k equal to

$$(P/P_k) \simeq \gamma^{-1} v_0^{-1} (j\omega - \gamma v_0)$$

Substituting (8.44), this becomes

$$(P/P_k) \simeq \{ \pm \omega_q / (\omega \mp \omega_q) \} \quad (8.55)$$

where the upper and lower signs apply to the fast and slow waves, respectively. Since ω is much greater than ω_q , (8.55) shows that power is mostly transmitted in the form of kinetic power, even in this extreme case with negligible displacement current.

In more general cases in which the electronic current is partially canceled by the displacement current, $\mathbf{H}_z$ becomes correspondingly smaller and the ratio of P to P_k will take some value between (8.49) and (8.55). We conclude from this that electromagnetic power is always small compared to the kinetic power if $\omega \gg \omega_p > \omega_q$, or equivalently, if ϱ_0 is small. When this is the case, the effect of the waveguide walls on the space charge waves is expected to be small, whether or not the walls are close to the electron beam. The majority of the transmission power is confined in the electron beam in the form of the kinetic power, and the electromagnetic field, which receives the constraint from the walls, only carries a small portion of the total power.

As we discussed earlier, there are a number of solutions for (8.35); and for each eigenfunction E_z , there are a pair of waves having propagation constants in the form of (8.44). Since E_z is different for different pairs, some pairs are strongly excited through external circuits while the others are not. In the case of a thin electron beam, the excitation of only one pair can be

made strong, and these generally have almost constant E_z over the beam cross section. For simplicity, let us concentrate on such a pair and neglect the effects of all other pairs. Since the current density of each wave is also almost constant over the beam cross section, the electronic current is given by the current density times the cross sectional area S_0 . Let J_+ and J_- be the electronic currents of the fast and slow waves at $z=0$. Then the total current at arbitrary z is given by

$$J(z) = J_+ e^{-\gamma_+ z} + J_- e^{-\gamma_- z} = (J_+ e^{j\beta_q z} + J_- e^{-j\beta_q z}) e^{-j\beta_e z} \quad (8.56)$$

where

$$\beta_q = (\omega_q/v_0) \quad (8.57)$$

On the other hand, from (8.23), (8.29), (8.31), and (8.44), the kinetic voltage is calculated to be

$$V_k(z) = Z_0 (J_+ e^{j\beta_q z} - J_- e^{-j\beta_q z}) e^{-j\beta_e z} \quad (8.58)$$

where

$$Z_0 = -\frac{m v_0 \omega_q}{e \rho_0 \omega} \frac{1}{S_0} = 2 \frac{V_0 \beta_q}{J_0 \beta_e} \quad (8.59)$$

The stationary component of the electron beam current is expressed by J_0 , and V_0 is the voltage applied to the electrons before they enter the region under consideration; hence, $\frac{1}{2} m v_0^2 = e V_0$.

The factor $\exp(-j\beta_e z)$ in (8.56) and (8.58) shows that the wave pattern as a whole moves in the positive z direction with the electron velocity v_0 . Except for this factor, $V_k(z)$ and $J(z)$ have functional forms similar to the voltage and current along a conventional transmission line having a characteristic impedance Z_0 and phase constant β_q . If we define the reflection coefficient r by

$$r = -(J_-/J_+) e^{-2j\beta_q z} \quad (8.60)$$

the impedance at z is given by

$$Z = \frac{V_k(z)}{J(z)} = Z_0 \frac{1+r}{1-r} \quad (8.61)$$

and the relation between r and Z can be studied using the Smith chart. Since

$$|V_k(z)| = Z_0 |J_+| |1+r| \quad (8.62)$$

$$|J(z)| = |J_+| |1-r| \quad (8.63)$$

we can easily see on the Smith chart how the magnitudes of $V_k(z)$ and $J(z)$ vary with z .

Let us consider the velocity of energy carried by each wave at angular frequency ω . Since the system is lossless, the velocity of energy is equal to the group velocity

$$v_g = \frac{d\omega}{d\beta} = \left(\frac{d\beta}{d\omega} \right)^{-1} = \left\{ \frac{1}{v_0} \left(1 \mp \frac{d\omega_q}{d\omega} \right) \right\}^{-1} \simeq v_0 \left(1 \pm \frac{d\omega_q}{d\omega} \right) \quad (8.64)$$

The upper and lower signs apply to the fast and slow waves, respectively. Usually, ω_q is a slowly varying function of ω , and the energy propagates with a velocity close to v_0 and in the same direction as the electron motion. The small discrepancy between the electron and energy velocities exists because a small portion of energy is transmitted in the form of electromagnetic power, as discussed previously. Note that v_g is equal to v_0 for the one-dimensional model in which the electromagnetic power disappears.

The kinetic powers of the fast and slow waves are calculated from appropriate terms in (8.56) and (8.58) to be

$$P_k = Z_0 |J_+|^2 \quad \text{and} \quad P_k = -Z_0 |J_-|^2$$

respectively. In the conventional transmission line, a negative power appears when the wave propagates in the negative z direction. The slow wave has, however, a negative power while propagating in the positive z direction. The following interpretation is usually given to explain the origin of the

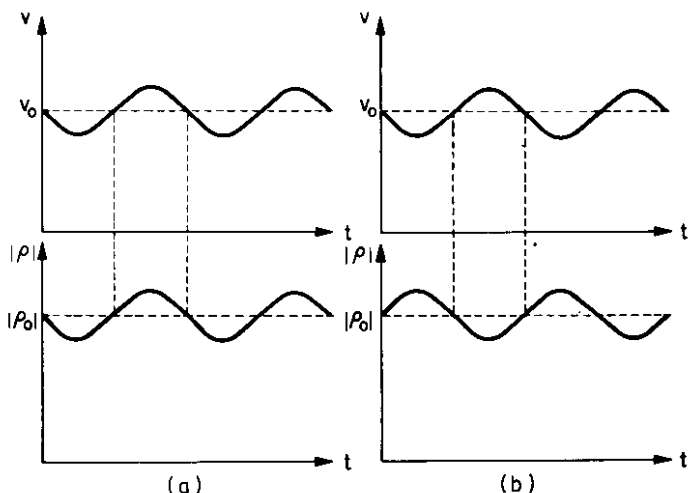


Fig. 8.2. The relation between electron velocity v and charge density $|\rho|$: (a) Fast wave; (b) Slow wave. (Note that ρ is negative.)

negative power. We obtain from (8.29), (8.30), and (8.44),

$$(v/\varrho) = (v_0/\varrho_0) \{ \pm \omega_q / (\omega \mp \omega_q) \} \quad (8.65)$$

where the upper signs apply to the fast wave and the lower signs to the slow wave. Equation (8.65) shows that if we observe the electron beam at a fixed point z , the electron velocity and the charge density vary with time, as illustrated in Fig. 8.2. For the fast wave, the velocity increases when the magnitude of the charge density, and hence the electron density, increases. Since most of the electrons increase their kinetic energy in this way, the kinetic energy of the electrons on the average increases when the fast wave is excited. For the slow wave, the electron density and the velocity are out of phase and consequently the average kinetic energy decreases when the slow wave is excited. This decrease of kinetic energy compared to the unexcited state propagates with the electrons in the positive z direction and is referred to as the negative kinetic power of the slow wave.

8.3 Gap Interactions

Let us consider two closely spaced planar grids which are capable of exciting a pair of space-charge waves for which E_z is almost constant over the electron beam cross section as discussed in the previous section. The planes of the grids are perpendicular to the beam axis, and the gap between them is narrow compared to $2\pi/\beta_e$. Furthermore, it is assumed that the grid wires are sufficiently thin and the meshes sufficiently large that no electrons are intercepted, and yet a perfect electro-static shield is provided by the grids. Under these assumptions, let us calculate the electronic current and kinetic voltage at the exit from the gap, assuming their values at the entrance are given. Let $z = -l$ and $z = 0$ be the locations of the first and second grids and V be the microwave voltage applied to the gap, as shown in Fig. 8.3. Since all the equations are linearized under the small signal assumption and the principle of superposition holds, the following two cases can be discussed separately:

- (i) The alternating voltage V is applied to the gap while keeping the current density $i_z(-l)$ and alternating electron velocity $v(-l)$ at the entrance equal to zero;
- (ii) V is equal to zero, but $i_z(-l)$ and $v(-l)$ have given values.

The superposition of these two cases gives the desired solution.

In case (i), the applied gap voltage V either accelerates or decelerates the

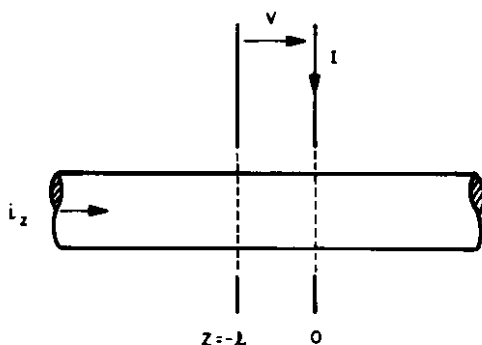


Fig. 8.3. The relation between V , I , and i_z at the electrode gap.

electrons depending on the phase of the voltage. The accelerated electrons tend to catch up the decelerated ones resulting in a modulation of charge density which in turn produces some E_z in addition to the original gap field. If the gap is narrow, as assumed here, the variation of charge density within the gap is negligible since the decelerated electrons leave the gap before the accelerated ones catch up. Under this assumption, the field E_z is equal to the applied voltage divided by the gap length l . The equation of motion for each electron is given by

$$\frac{d}{dt} v(z, t) = -\frac{e}{m} \left(\frac{-V}{l} \right) \sqrt{2} \cos \omega t \quad (8.66)$$

Since the electron velocity is approximately equal to v_0 , the electron velocity at $z = 0$ is obtained by adding the initial velocity v_0 at $z = -l$ to the integral of the above quantity with respect to t from $t - l/v_0$ to t :

$$\begin{aligned} v(0, t) &\simeq v_0 + \int_{t-l/v_0}^t \frac{e}{m} \frac{V}{l} \sqrt{2} \cos \omega t \, dt \\ &= v_0 + \frac{e}{m} \frac{V}{v_0} \frac{\sin(\theta/2)}{(\theta/2)} \sqrt{2} \cos(\omega t - \frac{1}{2}\theta) \end{aligned} \quad (8.67)$$

where

$$\theta = (\omega/v_0) l = \beta_e l \quad (8.68)$$

The term θ is called the transit angle of the gap. If we replace $\sqrt{2} V \cos \omega t$ by the corresponding complex expression $V e^{j\omega t}$, the alternating component of the velocity can be expressed in the form

$$v(0) = \frac{e}{m} \frac{V}{v_0} \frac{\sin(\theta/2)}{(\theta/2)} e^{-j(\theta/2)}$$

From this, the kinetic voltage V_k defined by (8.23) is given by

$$V_k(0) = MV \quad (8.69)$$

where

$$M = -\frac{\sin(\theta/2)}{(\theta/2)} e^{-j(\theta/2)} \quad (8.70)$$

The term M is called the beam-coupling factor. When l is sufficiently small, M is equal to -1 which means that the kinetic voltage at the exit of a narrow gap is equal to the applied gap voltage, except for the opposite sign. The corresponding $i_z(0)$ can be calculated by integrating (8.15). Since the alternating component q of charge density is equal to zero from the discussion preceding (8.66), and $i_z(-l) = 0$ by hypothesis, we have

$$i_z(0) = 0 \quad (8.71)$$

hence the electronic current J at $z = 0$ is equal to zero. This completes the discussion of case (i).

In case (ii), the gap exerts no force on the motion of electrons. Because of the phase difference due to the transit time, we multiply $v(-l)$ and $i_z(-l)$ by $e^{-j\theta}$ to obtain $v(0)$ and $i_z(0)$, respectively:

$$v(0) = v(-l) e^{-j\theta} \quad (8.72)$$

$$i_z(0) = i_z(-l) e^{-j\theta} \quad (8.73)$$

Note that $\beta_q l$ is neglected compared to $\beta_e l = \theta$, as comparison with (8.56) or (8.58) may indicate.

In general, $v(0)$ and $i_z(0)$ are given by the superposition of the above two cases; in terms of the kinetic voltage and electronic current, their values are

$$V_k(0) = V_k(-l) e^{-j\theta} + MV \quad (8.74)$$

$$J(0) = J(-l) e^{-j\theta} \quad (8.75)$$

So far, we have described how the electron beam is modulated by the gap voltage. Let us next consider the effect of the electron beam on the current flowing into the gap through the external circuit. To do so, we use the power relation (8.20) with the surface S enclosing the gap space. The first term on the left-hand side is the electromagnetic power flowing out through the surface S . Since $\mathbf{E}_t$ is equal to zero over the grid surfaces, no contribution to the integral results from these surfaces. The power flowing in from the side surface is given by VI^* , where V is the gap voltage and I the current. The second term on the left-hand side expresses the net kinetic power.

Since $V_k(-l) J^*(-l)$ enters through the first grid and $V_k(0) J^*(0)$ leaves through the second, the left-hand side of (8.20) becomes

$$- \{VI^* + V_k(-l) J^*(-l) - V_k(0) J^*(0)\} = - \{VI^* - MVJ^*(-l) e^{j\theta}\} \quad (8.76)$$

where the negative sign appears in front of the bracket because of the direction of $\mathbf{n}$.

The first and second terms on the right-hand side of (8.20) express the stored electric and magnetic energies, respectively. We neglect the second term because the inductance of the gap electrodes are usually small. The stored electric energy is calculated as follows, first, we note that the transverse component of the electric field is negligible since the gap is narrow and the tangential components of the field at the grid surfaces are zero. Second, the stored electric energy multiplied by $j\omega$ is given by

$$j\omega \epsilon |V/l|^2 l S_g = j\omega C V V^* \quad (8.77)$$

where $l S_g$ is the volume of the gap space, C is the gap capacitance given by $\epsilon S_g/l$, and the z -component of the electric field is approximated by V/l .

Finally, the third term on the right-hand side of (8.20) will be shown to be negligible. The integrand is

$$\frac{m}{e} \varrho_0 v(z) v^*(z) = \varrho_0 \frac{V_k(z) V_k^*(z)}{v_0^2 (m/e)} = \frac{1}{2} \frac{i_0}{V_0} |V_k(z)|^2 \frac{1}{v_0} \quad (8.78)$$

where V_0 is the static voltage required to accelerate electrons to the velocity v_0 , and $i_0 = \varrho_0 v_0$. Integrating the right-hand side over the volume occupied by the electron beam in the gap and multiplying the result by $j\omega$, we have

$$j \frac{1}{2} \frac{J_0}{V_0} \langle V_k^2 \rangle \frac{\omega l}{v_0} = j \frac{\langle V_k^2 \rangle}{Z_0} \beta_q l \quad (8.79)$$

where $\langle V_k^2 \rangle$ is the average value of $|V_k(z)|^2$ from $z = -l$ to 0 and Z_0 is defined by (8.59). Usually $\langle V_k^2 \rangle / Z_0$ is of the order of $V_k J$ and $\beta_q l \ll 1$, therefore the contribution from this term is negligibly small compared to the terms appearing in (8.76).

Summarizing the above discussion, (8.20) becomes

$$-VI^* = -MVJ^*(-l) e^{j\theta} + j\omega C V V^*$$

Dividing by $-V$ and taking the complex conjugate, we have

$$I = M^* J(-l) e^{-j\theta} + j\omega C V = MJ(-l) + j\omega C V \quad (8.80)$$

The first term on the right-hand side represents the current induced by the electron beam while the second term expresses the current into the gap capacitance. The sum of these two currents make up the net current flowing into the gap through the external circuit.

The above discussion applies to an ideal gap with two parallel grids, in practice however, a gridless gap is often utilized to avoid the interception of electrons as shown in Fig. 8.4. In such a case, the electric field at the electron

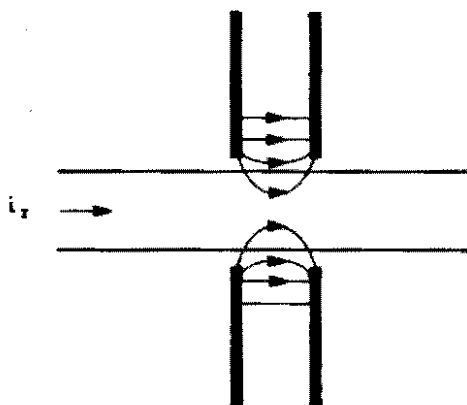


Fig. 8.4. Gridless electrode gap.

beam is weaker than the gap voltage divided by the gap distance, and the beam coupling factor M becomes smaller than that given by (8.70). To take this effect into account, M given by (8.70) is reduced by the factor $\eta < 1$.

Let us assume that the transit angle of the gap is very small so that

$$e^{-j\theta} \simeq 1, \quad \frac{\sin(\theta/2)}{(\theta/2)} \simeq 1 \quad (8.81)$$

then (8.74), (8.75), and (8.80) reduce to

$$V_k(0) = V_k(-l) - \eta V, \quad J(0) = J(-l), \quad I = -\eta J(-l) + j\omega CV \quad (8.82)$$

An equivalent circuit representing (8.82) is shown in Fig. 8.5. Note that the purpose of this equivalent circuit is simply to represent the relations (8.82), and it is not to be used for connecting impedances which fix the ratio between V_k and J at the ports corresponding to the electron beam.

In the remainder of this section, we shall discuss the principle of the

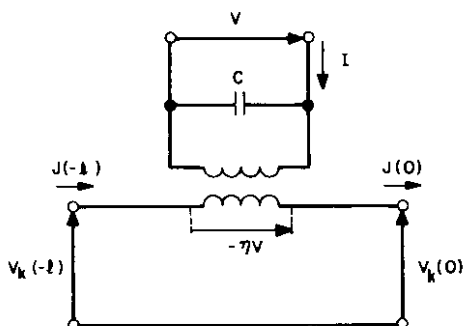


Fig. 8.5. Equivalent circuit of an electrode gap.

klystron as an example of utilizing (8.82). In its simplest form the klystron has two gaps placed one quarter of a plasma wavelength $\lambda_q = 2\pi/\beta_q$ apart. The first gap excites the space-charge waves while the second one receives power from the electron beam. Since the gap capacitance C is usually too large to apply the gap voltage or to extract power effectively, the capacitance effect is canceled by the inductance of a resonant cavity. Taking a proper reference plane, the equivalent circuit of the cavity including the gap should look like the one shown in Fig. 8.6, where G is the conductance representing the cavity losses. The parallel resonant circuit is dual to the series resonant circuit discussed in Section 4.3. In the following discussion, subscripts 1 and 2 will be used to express the quantities related to the first and second gaps, respectively.

In order to apply a voltage V_1 to the first gap, a net power $G_1|V_1|^2$ is

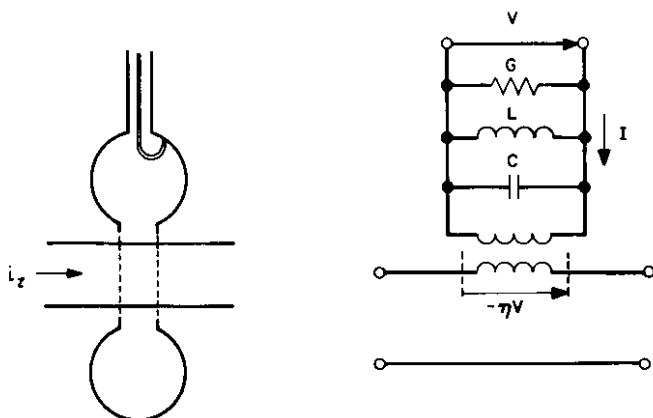


Fig. 8.6. Electrode gap with a resonant cavity and its equivalent circuit.

necessary. Since the available power P_a from the generator times the power transmission coefficient gives the net power, we have

$$P_a = \frac{G_1 |V_1|^2}{\{4G_1 G_g / (G_1 + G_g)^2\}} = \frac{|V_1|^2}{4G_g} (G_1 + G_g)^2 \quad (8.83)$$

where the cavity is assumed to resonate at the signal frequency, and G_g is the generator conductance as seen from the gap. If the electronic current J and kinetic voltage V_k are both equal to zero at the entrance to the first gap, V_k becomes $-\eta_1 V_1$, and J remains zero when leaving the gap. Taking the origin of the z -axis at the exit of the first gap, we have from (8.56) and (8.58)

$$J_+ + J_- = 0, \quad Z_0(J_+ - J_-) = -\eta_1 V_1$$

from which J_+ and J_- are obtained:

$$J_+ = -\frac{1}{2} Z_0^{-1} \eta_1 V_1, \quad J_- = \frac{1}{2} Z_0^{-1} \eta_1 V_1 \quad (8.84)$$

Thus, fast and slow waves having the same magnitude but opposite signs are excited by the first gap. As we move along the electron beam toward the second gap, from (8.62) and (8.63) the magnitudes of $V_k(z)$ and $J(z)$ change as shown in Fig. 8.7. At the entrance into the second gap $z_2 = \lambda_q/4$, $r = -1$, and

$$|V_k(z_2)| = 0, \quad |J(z_2)| = 2|J_+| \quad (8.85)$$

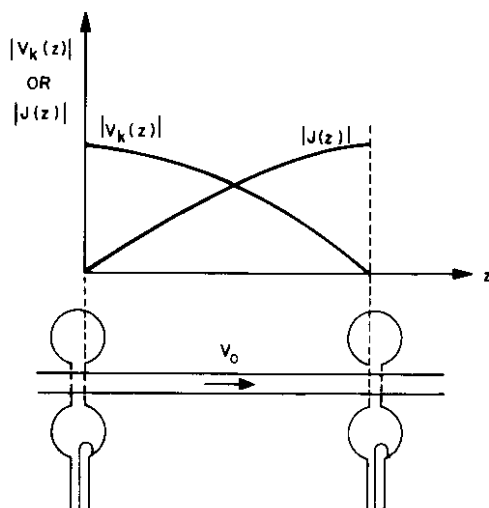


Fig. 8.7. Kinetic voltage and electronic current versus distance z .

Applying (8.82) to the second gap, we obtain

$$V_k(z_2 + l_2) = -\eta_2 V_2, \quad J(z_2 + l_2) = J(z_2), \quad I_2 = -\eta_2 J(z_2) + j\omega C V_2 \quad (8.86)$$

Let us assume that the capacitive current $j\omega C V_2$ is canceled by the inductive current in the cavity, then the total current flowing through the external load G_L plus the cavity loss conductance G_2 becomes $\eta_2 J(z_2)$, and the gap voltage V_2 is given by

$$V_2 = (G_2 + G_L)^{-1} \eta_2 J(z_2) \quad (8.87)$$

Using (8.84), (8.85), and (8.87), the actual power to the load is calculated to be

$$P_L = G_L |V_2|^2 = G_L (G_2 + G_L)^{-2} \eta_2^2 Z_0^{-2} \eta_1^2 |V_1|^2 \quad (8.88)$$

The transducer gain is given by

$$(P_L/P_a) = \eta_1^2 \eta_2^2 Z_0^{-2} 4G_g G_L (G_1 + G_g)^{-2} (G_2 + G_L)^{-2} \quad (8.89)$$

When $G_1 = G_g$ and $G_2 = G_L$, the maximum gain is obtained:

$$(P_L/P_a) = \eta_1^2 \eta_2^2 Z_0^{-2} (4G_1 G_2)^{-1}$$

The kinetic power at the entrance to the second gap is equal to zero since the kinetic voltage is zero. From (8.86) and (8.87), the outgoing kinetic power from the exit is given by

$$-\eta_2 V_2 J^*(z_2) = -(G_2 + G_L)^{-1} \eta_2^2 |J(z_2)|^2$$

which indicates that the electron beam receives this amount of negative kinetic power from the gap. On the other hand, from (8.87) the total power into the load G_L and G_2 is given by

$$(G_2 + G_L) |V_2|^2 = (G_2 + G_L)^{-1} \eta_2^2 |J(z_2)|^2$$

As expected, the power the electron beam receives plus the power to G_L and G_2 is equal to zero. In other words, the gap receives power from the electron beam and delivers it to G_L and G_2 .

In order to increase the maximum gain, G_1 and G_2 must be decreased which is obvious since G_1 and G_2 express the cavity losses. On the other hand, some explanation may be necessary as to why G_g and G_L have to be decreased in order to increase the transducer gain when G_1 and G_2 can be neglected. Since the electronic current J is equal to zero at the exit from the first gap, no kinetic power is added to the electron beam by the gap. This

means that no power is required to excite the space-charge waves in the first gap although a large gap voltage is necessary for a large gain. To increase the gap voltage V_1 , keeping P_a constant, G_g has to be decreased. For the second gap, the active component of the current I_2 is solely determined by the electronic current $J(z_2)$. In order to increase the output power under the constant current condition, the voltage V_2 has to be increased by decreasing G_L . Hence, G_g and G_L must be small for a large gain.

Summarizing the above discussion, a klystron amplifies a signal because a negligible amount of power is required to excite space-charge waves in a narrow gap while a second gap one quarter of a plasma wavelength away provides a constant current source from which useful power can be extracted.

8.4 Continuous Interaction with a Slow Wave Structure

It is possible to achieve amplification utilizing the continuous interaction between an electron beam and an electromagnetic wave on a slow wave structure when their velocities are nearly equal. An electron tube based on this principle is called a traveling wave tube. The slow wave structure for a traveling wave tube is often realized by winding a thin wire into a helix with its diameter sufficiently small compared to a wavelength in free space. The electromagnetic wave propagating along the helix tends to follow the wound wire, and its axial phase velocity can be made approximately equal to the electron velocity when the pitch of the helix is properly chosen. In addition to helices, many periodic structures of the type discussed in Sec. 6.4 can be utilized as slow wave structures in traveling wave tubes.

Let $a_1(z)$ represent the electromagnetic wave propagating along the slow wave structure in the positive z direction. By suitable normalization, $|a_1(z)|^2$ is made to express the transmission power. The interaction between the electron beam and the slow wave structure takes place through the z -component of the electric field with a proper phase constant. Since E_z is 90° out of phase from E_t , as (3.61) indicates, if the phase of $a_1(z)$ is chosen to be in phase with the corresponding E_t , the z -component of the electric field at the electron beam in the slow wave structure becomes

$$E_z e^{-j\beta_s z} = j\beta_s Z_k^{1/2} a_1(z) \quad (8.90)$$

where β_s represents the phase constant of the proper field component (i.e., $\beta_s \simeq \beta_e$) and $Z_k^{1/2}$ is a proportional constant. Note that Z_k has the dimension of impedance.

Let us next express the space-charge waves, along the electron beam by

$$a_2(z) = Z_0^{1/2} J_+ e^{j(\beta_a - \beta_e)z}, \quad a_3(z) = Z_0^{1/2} J_- e^{-j(\beta_a + \beta_e)z} \quad (8.91)$$

then the transmission powers in the positive z direction due to the space-charge waves are given by $|a_2(z)|^2$ and $-|a_3(z)|^2$, respectively, since the electromagnetic power of each wave is negligible compared to its kinetic power.

Assuming a small coupling between the electron beam and the slow wave structure, we now apply the theory of coupled modes described in Section 6.1. To do so, however, we must obtain the coupling coefficients C_{21} and C_{31} . Suppose that $a_2(z)$ and $a_3(z)$ are zero at $z = 0$; let us investigate their values at $z = \Delta z$ for a given $a_1(0)$. From (8.90), E_z is approximately given by $j\beta_s Z_k^{1/2} a_1(0)$ over the narrow region from $z = 0$ to Δz . If we set $l = \Delta z$ and $V = -E_z \Delta z = -j\beta_s Z_k^{1/2} a_1(0) \Delta z$, then the problem reduces to that of the gap interaction discussed in Section 8.3. In the present case, $\Delta z = l$ is an infinitesimal, and the beam coupling factor M becomes -1 . Thus, the kinetic voltage at $z = \Delta z$ is given by

$$V_k(\Delta z) = -V = j\beta_s Z_k^{1/2} a_1(0) \Delta z \quad (8.92)$$

The electronic current remains the same:

$$J(\Delta z) = 0 \quad (8.93)$$

Combining (8.91), (8.92), and (8.93), we have

$$a_2(\Delta z) - a_3(\Delta z) = j\beta_s (Z_k/Z_0)^{1/2} a_1(0) \Delta z, \quad a_2(\Delta z) + a_3(\Delta z) = 0 \quad (8.94)$$

from which $a_2(\Delta z)$ and $a_3(\Delta z)$ are calculated to be

$$a_2(\Delta z) = \frac{1}{2} j\beta_s (Z_k/Z_0)^{1/2} a_1(0) \Delta z \quad (8.95)$$

$$a_3(\Delta z) = -\frac{1}{2} j\beta_s (Z_k/Z_0)^{1/2} a_1(0) \Delta z \quad (8.96)$$

Comparing these equations with (6.4) and remembering that $a_2(0) = a_3(0) = 0$, C_{21} and C_{31} are found to be

$$C_{21} = -C_{12}^* = -\frac{1}{2} j\beta_s (Z_k/Z_0)^{1/2}, \quad C_{31} = C_{13}^* = \frac{1}{2} j\beta_s (Z_k/Z_0)^{1/2} \quad (8.97)$$

The coupling coefficient C_{23} between the fast and slow space-charge waves will be neglected since it becomes a higher order infinitesimal when we impose the condition of small coupling between the electron beam and the slow wave structure.

For simplicity, let us first consider the interaction between two waves. When $a_1(z)$ interacts with $a_2(z)$, the two waves exchange transmission power

back and forth as we discussed in connection with Fig. 6.3. On the other hand, if $a_1(z)$ and $a_3(z)$ have similar phase velocities and interact with each other strongly, growing waves will result as we discussed in connection with Fig. 6.4. Let

$$\beta_0 + \Delta\beta_0 = \beta_s, \quad \beta_0 - \Delta\beta_0 = \beta_q + \beta_e \quad (8.98)$$

and assume that $a_1(z) = A_0$ and $a_3(z) = 0$ at $z = 0$, the input of the coupling region, then

$$|a_1(z)|^2 = A_0^2 \{1 + (|C_{13}|/\alpha)^2 \sinh^2 \alpha z\} \quad (8.99)$$

$$-|a_3(z)|^2 = -A_0^2 (|C_{13}|/\alpha)^2 \sinh^2 \alpha z \quad (8.100)$$

where

$$\alpha = \{|C_{13}|^2 - (\Delta\beta_0)^2\}^{1/2} \quad (8.101)$$

and $|C_{13}|^2 > (\Delta\beta_0)^2$ is assumed. If the two waves are separated at $z = L$ and the electromagnetic wave along the slow wave structure is fed into a matched load, the transducer gain is given by

$$|a_1(L)|^2/|a_1(0)|^2 = 1 + (|C_{13}|/\alpha)^2 \sinh^2 \alpha L \quad (8.102)$$

provided that the input is also matched. This explains the amplification mechanism of traveling wave tubes.

In practice, β_q and $|C_{13}|$ often have the same order of magnitude, in such a case the interactions among three waves, $a_1(z)$, $a_2(z)$, and $a_3(z)$ must be considered. The eigenvalue problem for the system is given by

$$\begin{bmatrix} j\beta_s & -C_{13}^* & C_{13} \\ C_{13} & j(\beta_e - \beta_q) & 0 \\ C_{13}^* & 0 & j(\beta_e + \beta_q) \end{bmatrix} \mathbf{x} = \gamma \mathbf{x} \quad (8.103)$$

The eigenvalues are obtained from the solutions of the third order algebraic equation

$$\begin{vmatrix} j\beta_s - \gamma & -C_{13}^* & C_{13} \\ C_{13} & j(\beta_e - \beta_q) - \gamma & 0 \\ C_{13}^* & 0 & j(\beta_e + \beta_q) - \gamma \end{vmatrix} = 0 \quad (8.104)$$

If we set

$$\gamma = j\beta_e + \delta \quad (8.105)$$

Eq. (8.104) becomes

$$(\delta^2 + \beta_q^2) \{-\delta + j(\beta_s - \beta_e)\} + 2j\beta_q |C_{13}|^2 = 0 \quad (8.106)$$

Let α and γ be the real and imaginary parts of δ :

$$\delta = \alpha + jy \quad (8.107)$$

then (8.106) can be decomposed into real and imaginary parts as follows:

$$\begin{aligned} -\alpha[\alpha^2 - y^2 + \beta_q^2 - 2y\{y - (\beta_s - \beta_e)\}] &= 0 \\ (-y + \beta_s - \beta_e)(\alpha^2 - y^2 + \beta_q^2) - 2\alpha^2 y + 2\beta_q |C_{13}|^2 &= 0 \end{aligned} \quad (8.108)$$

By inspection, there are solutions satisfying

$$\alpha = 0, \quad \beta_s - \beta_e = 2\beta_q |C_{13}|^2 (y^2 - \beta_q^2)^{-1} + y \quad (8.109)$$

The waves corresponding to these solutions do not grow or decay with z since the real part of γ is equal to zero. If we assume nonzero α (8.108) gives

$$\alpha = \pm [2y\{y - (\beta_s - \beta_e)\} + y^2 - \beta_q^2]^{1/2} \quad (8.110)$$

$$\beta_s - \beta_e = 2y \pm \{\beta_q^2 + \beta_q |C_{13}|^2 y^{-1}\}^{1/2} \quad (8.111)$$

Suppose that the value of $\beta_q/|C_{13}|$ is fixed, then one value of $\beta_s - \beta_e$ can be calculated from (8.109) for each value of y . Similarly, (8.111) gives the other values of $\beta_s - \beta_e$, and (8.110) gives the corresponding α , provided that α and $\beta_s - \beta_e$, thus obtained, are all real. In this way, we can obtain all the possible values of α and $\beta_s - \beta_e$ for each y from which we can reconstruct the propagation constant γ as a multivalued function of $\beta_s - \beta_e$. As an example, the result of the calculation for a particular case in which $\beta_q = \sqrt{2}|C_{13}|$ is shown in Fig. 8.8. To facilitate a direct comparison with Fig. 6.2, the real and imaginary components of $\gamma/|C_{13}|$ are shown as functions of $\Delta\beta_0/|C_{13}|$. The interaction between $a_1(z)$ and $a_2(z)$ gives Fig. 6.2(a), and that between $a_1(z)$ and $a_3(z)$ gives Fig. 6.2(b), however, when $a_1(z)$ interacts with $a_2(z)$ and $a_3(z)$, as in the present case, the two patterns are intermingled and distorted as illustrated in Fig. 8.8.

Once the eigenvalues are obtained, the corresponding eigenvectors can easily be calculated from (8.103):

$$\mathbf{x}_i = \begin{bmatrix} 1 \\ -C_{13}/(-\delta_i - j\beta_q) \\ -C_{13}^*/(-\delta_i + j\beta_q) \end{bmatrix} \quad (i = 1, 2, 3) \quad (8.112)$$

Note that three eigenvectors are obtained corresponding to three different eigenvalues. Suppose that $\text{Re}\gamma_i = \alpha_i < 0$ and that $\mathbf{x}_1$ represents the corresponding growing wave, then the amplitude of $\mathbf{x}_1$ increases with z while the other waves either decay or propagate unchanged, hence the $\mathbf{x}_1$ component outweighs the others at the output port provided that the coupling region is sufficiently long.

Under the conditions of the previous paragraph the transducer gain can be calculated as follows: We assume that $a_1(z) = A_0$ at the input $z = 0$,

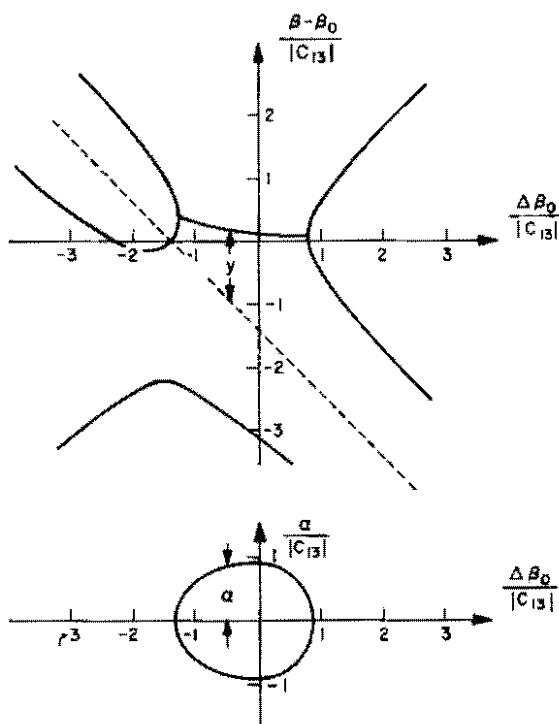


Fig. 8.8. $\beta - \beta_0$ and α versus $\Delta\beta_0 = \frac{1}{2}(\beta_s - \beta_e - \beta_q)$.

then the x_1 component at $z = 0$ is given by

$$\tilde{x}_1 \mathbf{P} \begin{bmatrix} A_0 \\ 0 \\ 0 \end{bmatrix} x_1$$

where

$$\tilde{x}_1 = \frac{\delta_1^2 + \beta_q^2}{3\delta_1^2 + \beta_q^2 - 2\delta_1 j(\beta_s - \beta_e)} \begin{bmatrix} 1 & -C_{13}^* & -C_{13} \\ \delta_1 + j\beta_q & \delta_1 - j\beta_q \end{bmatrix}$$

At the output where $z = L$, the magnitude of the x_1 component becomes $\exp(-\alpha_1 L)$ times as large, and hence the transducer gain becomes

$$\left| \frac{a_1(L)}{a_1(0)} \right|^2 \simeq \left| \frac{\delta_1^2 + \beta_q^2}{3\delta_1^2 + \beta_q^2 - 2\delta_1 j(\beta_s - \beta_e)} \right|^2 e^{-2\alpha_1 L} \quad (8.113)$$

8.5 Electron Beam Noise

Let us assume that the fast and slow waves are given by A_1 and A_2 , respectively, at a cross section a of the electron beam shown in Fig. 8.9, and that they become B_1 and B_2 at another cross section b situated beyond a lossless circuit. Let $\mathbf{a}$ and $\mathbf{b}$ be defined by

$$\mathbf{a} = \begin{bmatrix} A_1 \\ A_2 \end{bmatrix}, \quad \mathbf{b} = \begin{bmatrix} B_1 \\ B_2 \end{bmatrix} \quad (8.114)$$

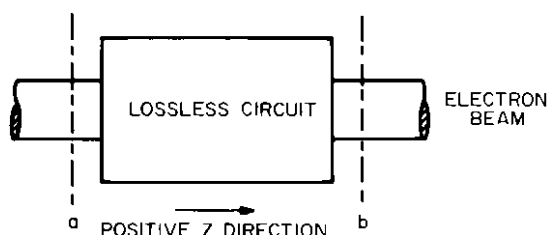


Fig. 8.9. Lossless transformation of space charge waves.

Under the small signal assumption, the system is linear and $\mathbf{b}$ must be expressible as a square matrix $\mathbf{K}$ of order 2 times $\mathbf{a}$:

$$\mathbf{b} = \mathbf{K}\mathbf{a} \quad (8.115)$$

The total transmission power in the positive z direction at a is given by

$$|A_1|^2 - |A_2|^2 = \mathbf{a}^+ \mathbf{P} \mathbf{a} \quad (8.116)$$

where

$$\mathbf{P} = \text{diag}[1 \quad -1] \quad (8.117)$$

The negative sign in front of $|A_2|^2$ comes from the fact that the slow wave carries a negative power. The total transmission power at b is given by $\mathbf{b}^+ \mathbf{P} \mathbf{b}$, and since the system is lossless, these two powers must be equal, i.e.,

$$\mathbf{a}^+ \mathbf{P} \mathbf{a} - \mathbf{b}^+ \mathbf{P} \mathbf{b} = 0$$

or equivalently,

$$\mathbf{a}^+ (\mathbf{P} - \mathbf{K}^+ \mathbf{P} \mathbf{K}) \mathbf{a} = 0 \quad (8.118)$$

Since a is arbitrary, we have

$$\mathbf{K}^+ \mathbf{P} \mathbf{K} = \mathbf{P} \quad (8.119)$$

which expresses the lossless condition for $\mathbf{K}$. Taking the determinant of

(8.119), we see that $\det \mathbf{K}$ is not equal to zero, and hence $\mathbf{K}^{-1}$ exists.

In order to discuss electron beam noise, let us consider the noise power matrix $\langle \mathbf{a}\mathbf{a}^+ \rangle$ where $\langle \rangle$ indicates the ensemble average. For example, $\langle A_1 A_1^* \rangle$, which is the 11 component of $\langle \mathbf{a}\mathbf{a}^+ \rangle$, expresses the noise power in the fast wave in a narrow bandwidth B while the 12 component $\langle A_1 A_2^* \rangle$ expresses the correlation between the noise components in the fast and slow waves. At the cross section b , the noise power matrix becomes $\langle \mathbf{b}\mathbf{b}^+ \rangle$, and because of (8.115), it can be written in the form

$$\langle \mathbf{b}\mathbf{b}^+ \rangle = \mathbf{K} \langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{K}^+ \quad (8.120)$$

Multiplying (8.120) by $\mathbf{P}$ from the right and using (8.119), we have

$$\langle \mathbf{b}\mathbf{b}^+ \rangle \mathbf{P} = \mathbf{K} \langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{K}^+ \mathbf{P} = \mathbf{K} \langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{P} \mathbf{K}^{-1} \quad (8.121)$$

Note that $\langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{P}$ undergoes the similarity transformation by $\mathbf{K}$ when the electron beam passes through the lossless circuit. Since the trace and determinant of a matrix are invariant to a similarity transformation, n_p and n_s defined through

$$\text{tr}(\langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{P}) = \langle |A_1|^2 \rangle - \langle |A_2|^2 \rangle \equiv n_p \quad (8.122)$$

$$\det(\langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{P}) = \langle |A_1 A_2^*|^2 \rangle - \langle |A_1|^2 \rangle \langle |A_2|^2 \rangle \equiv \frac{1}{4}(n_p^2 - n_s^2) \quad (8.123)$$

are invariant.

We shall now discuss the measurement of the noise parameters n_s and n_p . In terms of A_1 and A_2 , the electron beam current is given by

$$J(z) = (Z_0)^{-1/2} \{A_1 e^{i\beta_0 z} + A_2 e^{-i\beta_0 z}\} e^{-i\beta_0 z} \quad (8.124)$$

from which we have

$$\langle |J(z)|^2 \rangle = Z_0^{-1} \{ \langle |A_1|^2 \rangle + \langle |A_2|^2 \rangle + 2 \langle |A_1 A_2^*| \rangle \cos(2\beta_0 z + \varphi) \} \quad (8.125)$$

where φ is defined through

$$\tan \varphi = \{ \text{Im}(\langle A_1 A_2^* \rangle) / \text{Re}(\langle A_1 A_2^* \rangle) \} \quad (8.126)$$

If we move a gap coupled resonant cavity along the electron beam, as shown in Fig. 8.10(a), the noise output from the cavity is proportional to $\langle |J(z)|^2 \rangle$. Suppose that the electron beam direct current J_0 is reduced to a very small value, then the cavity output corresponds to the full shot noise $2eBJ_0 = \langle |J(z)|^2 \rangle$. Calibrating the cavity output using this shot noise, the absolute value of $\langle |J(z)|^2 \rangle$ under the operating condition can be obtained from the above measurement.

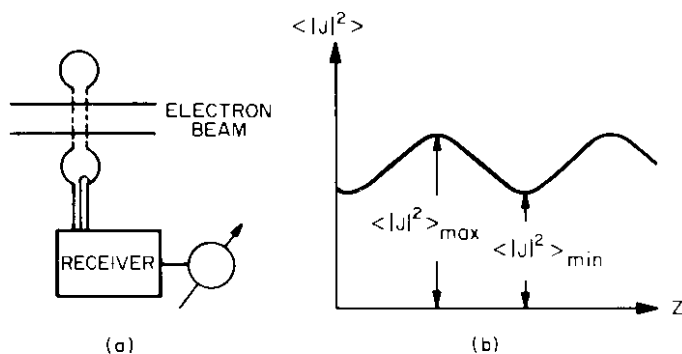


Fig. 8.10. Measurement of noise parameters n_s and n_p .

The product of the maximum and minimum values of $\langle |J(z)|^2 \rangle$ is calculated from (8.125) to be

$$\langle |J|^2 \rangle_{\max} \langle |J|^2 \rangle_{\min} = Z_0^{-2} \{ (\langle |A_1|^2 \rangle + \langle |A_2|^2 \rangle)^2 - 4 \langle |A_1 A_2^*|^2 \rangle \} \quad (8.127)$$

From (8.122) and (8.123) the terms inside the brackets on the right-hand side are equal to n_s^2 , so that

$$n_s^2 = Z_0^2 \langle |J|^2 \rangle_{\max} \langle |J|^2 \rangle_{\min} \quad (8.128)$$

The impedance Z_0 can be calculated from V_0 , J_0 , β_q , and β_e , and the value of $\langle |J|^2 \rangle_{\max} \langle |J|^2 \rangle_{\min}$, is obtainable from the measurement of $\langle |J(z)|^2 \rangle$. Thus, n_s can be determined from (8.128).

To obtain n_p , let us next take the ratio of the maximum and minimum values of $\langle |J(z)|^2 \rangle$:

$$\varrho^2 \equiv \frac{\langle |J|^2 \rangle_{\max}}{\langle |J|^2 \rangle_{\min}} = \frac{\langle |A_1|^2 \rangle + \langle |A_2|^2 \rangle + 2 \langle |A_1 A_2^*|^2 \rangle}{\langle |A_1|^2 \rangle + \langle |A_2|^2 \rangle - 2 \langle |A_1 A_2^*|^2 \rangle} \quad (8.129)$$

If there is no correlation between A_1 and A_2 , $\varrho^2 = 1$, and the cavity output remains the same regardless of its position along the beam. Only when A_1 and A_2 are correlated, $\langle |J(z)|^2 \rangle$ changes as shown in Fig. 8.10(b). From (8.122) and (8.123), n_p/n_s can be expressed in terms of ϱ , $\langle |A_1|^2 \rangle$ and $\langle |A_2|^2 \rangle$ as follows:

$$\frac{n_p}{n_s} = \frac{\varrho^2 + 1}{2\varrho} \frac{\langle |A_1|^2 \rangle - \langle |A_2|^2 \rangle}{\langle |A_1|^2 \rangle + \langle |A_2|^2 \rangle} \quad (8.130)$$

This shows that n_p can be calculated from the previously measured n_s and

ϱ if the relative magnitudes of $\langle |A_1|^2 \rangle$ and $\langle |A_2|^2 \rangle$ are obtained. The terms $\langle |A_1|^2 \rangle$ and $-\langle |A_2|^2 \rangle$ express the fast and slow wave noise powers, respectively. Suppose that two cavities with each gap coupled to the electron beam are separated by a certain distance and that their outputs are combined through appropriate phase shifters, then the resultant output can be made to couple either with the fast wave only or with the slow wave; in this way the relative magnitudes of $|A_1|^2$ and $|A_2|^2$ can be obtained. Thus, both n_s and n_p are measurable with movable cavities coupled to the electron beam.

Let us finally consider the best amplifier noise performance obtainable from an electron beam with given n_s and n_p . To do so, we first try to obtain a canonical form of the noisy electron beam. The matrix $\langle \mathbf{a}\mathbf{a}^+ \rangle$ is self-adjoint and can be either positive-definite or positive-semidefinite; we shall restrict ourselves to the case in which $\langle \mathbf{a}\mathbf{a}^+ \rangle$ is positive-definite, then we have two self-adjoint matrices $\mathbf{P}$ and $\langle \mathbf{a}\mathbf{a}^+ \rangle$ one of which is positive-definite. From the discussion leading to (5.60) and (5.61), they can be simultaneously diagonalized by a matrix $\mathbf{H}$ as follows.

$$\mathbf{H}^+ \langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{H} = \mathbf{I} \quad (8.131)$$

$$\mathbf{H}^+ \mathbf{P} \mathbf{H} = \begin{bmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{bmatrix} \quad (8.132)$$

where both λ_1 and λ_2 are real. Taking the determinant of the left-hand side in (8.132), we obtain

$$\det \mathbf{H}^+ \det \mathbf{P} \det \mathbf{H} = \det \mathbf{P} |\det \mathbf{H}|^2$$

which is negative. Consequently, the determinant of the right-hand side, $\lambda_1 \lambda_2$, must be negative, in other words, λ_1 and λ_2 have opposite signs. If the first diagonal component of (8.132) happens to be negative for a certain $\mathbf{H}$, use $\mathbf{H}\mathbf{L}$ in place of $\mathbf{H}$, where $\mathbf{L}$ is given by

$$\mathbf{L} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix}$$

then we have

$$\mathbf{L}\mathbf{H}^+ \langle \mathbf{a}\mathbf{a}^+ \rangle \mathbf{H}\mathbf{L} = \mathbf{L}\mathbf{I}\mathbf{L} = \mathbf{I}$$

$$\mathbf{L}\mathbf{H}^+ \mathbf{P} \mathbf{H}\mathbf{L} = \mathbf{L} \begin{bmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{bmatrix} \mathbf{L} = \begin{bmatrix} \lambda_2 & 0 \\ 0 & \lambda_1 \end{bmatrix}$$

and the first component becomes positive. Without loss of generality, we can therefore assume that λ_1 is positive in (8.132) and that λ_2 is negative.

Let Γ be defined by

$$\Gamma = \begin{bmatrix} (\lambda_1)^{-1/2} & 0 \\ 0 & (-\lambda_2)^{-1/2} \end{bmatrix} \quad (8.133)$$

then (8.132) can be written in the form

$$\mathbf{H}^+ \mathbf{P} \mathbf{H} = \Gamma^{-1} \mathbf{P} \Gamma^{-1}$$

which is equivalent to

$$\Gamma \mathbf{H}^+ \mathbf{P} \mathbf{H} \Gamma = \mathbf{P} \quad (8.134)$$

On the other hand, from (8.131) we have

$$\Gamma \mathbf{H}^+ \langle \mathbf{a} \mathbf{a}^+ \rangle \mathbf{H} \Gamma = \begin{bmatrix} \lambda_1^{-1} & 0 \\ 0 & -\lambda_2^{-1} \end{bmatrix} \quad (8.135)$$

If we set

$$\mathbf{K} = \Gamma \mathbf{H}^+ \quad (8.136)$$

then (8.134) and (8.135) become

$$\mathbf{K} \mathbf{P} \mathbf{K}^+ = \mathbf{P} \quad (8.137)$$

$$\mathbf{K} \langle \mathbf{a} \mathbf{a}^+ \rangle \mathbf{K}^+ = \begin{bmatrix} \lambda_1^{-1} & 0 \\ 0 & -\lambda_2^{-1} \end{bmatrix} \quad (8.138)$$

Following the derivation of (7.31), (8.119) shows that $\mathbf{K}$ satisfying (8.137) represents a lossless transformation. Multiplying (8.138) by $\mathbf{P}$ from the right and comparing the result with (8.121), we have

$$\text{tr} \left(\begin{bmatrix} \lambda_1^{-1} & 0 \\ 0 & -\lambda_2^{-1} \end{bmatrix} \mathbf{P} \right) = \lambda_1^{-1} + \lambda_2^{-1} = n_p \quad (8.139)$$

$$\det \left(\begin{bmatrix} \lambda_1^{-1} & 0 \\ 0 & -\lambda_2^{-1} \end{bmatrix} \mathbf{P} \right) = \lambda_1^{-1} \lambda_2^{-1} = \frac{1}{4} (n_p^2 - n_s^2) \quad (8.140)$$

where use is made of the fact that the trace and determinant of $\langle \mathbf{a} \mathbf{a}^+ \rangle \mathbf{P}$ are invariant to lossless transformations. From (8.139) and (8.140), we obtain

$$\lambda_1^{-1} = \frac{1}{2} (n_s + n_p), \quad -\lambda_2^{-1} = \frac{1}{2} (n_s - n_p)$$

Let $\langle \mathbf{a}' \mathbf{a}'^+ \rangle$ be the noise power matrix of the electron beam after the lossless transformation by $\mathbf{K}$, then we have

$$\langle \mathbf{a}' \mathbf{a}'^+ \rangle = \begin{bmatrix} \frac{1}{2} (n_s + n_p) & 0 \\ 0 & \frac{1}{2} (n_s - n_p) \end{bmatrix} \quad (8.141)$$

The matrix $\mathbf{K}$ transforms $\mathbf{a}$ to $\mathbf{a}'$ and $\mathbf{K}^{-1}$ transforms $\mathbf{a}'$ back to $\mathbf{a}$. Note that $\mathbf{K}^{-1}$ also represents a lossless transformation since $(\mathbf{K}^{-1})^+ \mathbf{P} \mathbf{K}^{-1} = \mathbf{P}$ from

(8.119). We conclude from the above discussion that an electron beam with $\langle aa^+ \rangle$ can be considered as the result of the lossless transformation $\mathbf{K}^{-1}$ applied to a similar electron beam with $\langle a'a'^+ \rangle$. The electron beam with $\langle a'a'^+ \rangle$ corresponds to the canonical form of linear noisy networks discussed in Section 7.1, and $\mathbf{K}$ corresponds to the canonical transformation.

We are now in a position to calculate the optimum noise measure of the amplifier shown in Fig. 8.11. The power waves at the input are given by a_1

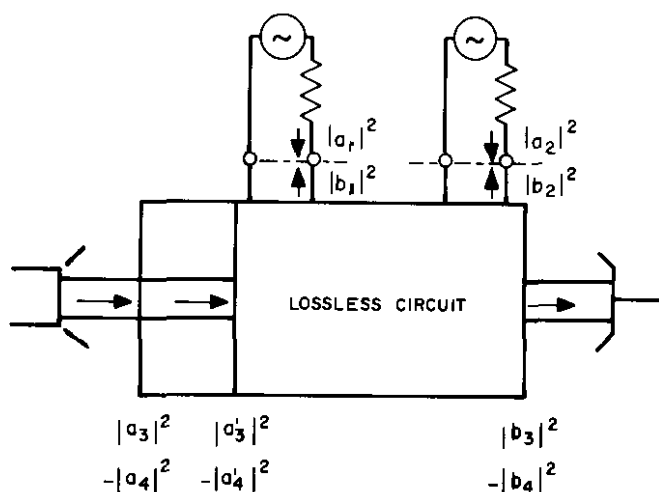


Fig. 8.11. An electron beam amplifier.

and b_1 , and at the output by a_2 and b_2 . The generator and load are assumed to have impedances with positive real parts. Let a_3 and a_4 be the noise components in the fast and slow waves at the entrance to the coupling region, and let b_3 and b_4 be those at the exit. Since a_3 and a_4 can be transformed by a lossless circuit to uncorrelated waves a_3' and a_4' , as we discussed above, we can use a_3' and a_4' instead of a_3 and a_4 to investigate the effect of the most general lossless network.

Let us define a scattering matrix $\mathbf{S}$ through

$$\begin{bmatrix} b_1 \\ b_2 \\ b_3 \\ b_4 \end{bmatrix} = \mathbf{S} \begin{bmatrix} a_1 \\ a_2 \\ a_3' \\ a_4' \end{bmatrix} \quad (8.142)$$

then the gain and noise measure of the amplifier are given by

$$G = |S_{21}|^2 \quad (8.143)$$

$$M = \frac{|S_{22}|^2 |a_2|^2 + |S_{23}|^2 |a_3'|^2 + |S_{24}|^2 |a_4'|^2}{(|S_{21}|^2 - 1) |a_1|^2} \quad (8.144)$$

where

$$|a_1|^2 = kT_i B, \quad |a_2|^2 = kT_L B, \quad |a_3'|^2 = \frac{1}{2}(n_s + n_p), \quad |a_4'|^2 = \frac{1}{2}(n_s - n_p) \quad (8.145)$$

Since $\mathbf{S}$ represents a lossless network, it satisfies

$$\mathbf{S}^+ \mathbf{P} \mathbf{S} = \mathbf{P}$$

where $\mathbf{P} = \text{diag}[1 \ 1 \ 1 \ -1]$. From this, following the derivation of (7.32), we obtain

$$|S_{21}|^2 + |S_{22}|^2 + |S_{23}|^2 - |S_{24}|^2 = 1$$

Substituting this relation into (8.144), we have

$$M = \frac{|S_{22}|^2 |a_2|^2 + |S_{23}|^2 |a_3'|^2 + |S_{24}|^2 |a_4'|^2}{(-|S_{22}|^2 - |S_{23}|^2 + |S_{24}|^2) |a_1|^2} \quad (8.146)$$

from which the optimum value of M is calculated to be $|a_4'|^2/|a_1|^2$, or equivalently

$$M_{\text{opt}} = \{(n_s - n_p)/2kT_i B\} \quad (8.147)$$

where use is made of (8.145). An argument similar to the one given in Section 7.2 shows that it is impossible to achieve a positive M smaller than M_{opt} with any passive network.

In the above discussion, we used a small signal assumption and considered only one pair of space-charge waves. In practice, the small signal assumption will not hold in the vicinity of the cathode where v_0 is small, hence there is a possibility of improving the optimum noise measure of an electron beam by modifying the potential profile near the cathode. Consequently, a major effort has been concentrated on the study of cathodes and their vicinity in order to obtain low noise electron beam amplifiers.

PROBLEMS

- 8.1 Calculate the plasma frequency of an electron beam with diameter 1 mm, direct current $J_0 = 10$ mA, and accelerating voltage $V_0 = 1000$ V.
- 8.2 Prove that the ratio between the electromagnetic and kinetic powers of a uniform

electron beam filling the inside of a waveguide is given by

$$\frac{P}{P_k} \simeq \frac{\pm \omega_q}{\omega \mp \omega_q} \frac{\omega_p^2 - \omega_q^2}{\omega_p^2}$$

where the upper signs apply to the fast wave and the lower signs to the slow wave.

- 8.3 Suppose that one additional cavity is gap-coupled to the electron beam midway between the input and output cavities of a klystron; calculate the gain as a function of the distance between adjacent cavities.
- 8.4 Solve the eigenvalue problem (8.35) for a uniform electron beam with radius a located coaxially inside a circular waveguide with radius b ($b > a$). Assume rotational symmetry of the modes to simplify the calculation.

(Hint: The propagation constant γ of each mode can be obtained from (8.43) with k_n satisfying

$$\frac{k_n J_0'(k_n a)}{J_0(k_n a)} = \frac{\beta_e \{I_0'(\beta_e a) K_0(\beta_e b) - K_0'(\beta_e a) I_0(\beta_e b)\}}{I_0(\beta_e a) K_0(\beta_e b) - K_0(\beta_e a) I_0(\beta_e b)}$$

where $I_0(\beta_e r)$ and $K_0(\beta_e r)$ are hyperbolic Bessel functions and use has been made of $\omega^2 \epsilon \mu + \gamma^2 \simeq -\beta_e^2$. The hyperbolic Bessel functions are two independent solutions of the differential equation $(d^2 u/dr^2) + (1/r)(du/dr) - \beta_e^2 u = 0$.)

- 8.5 Calculate the gain of a traveling wave tube without assuming the predominance of the growing wave at the output end of the coupling region.

CHAPTER 9

OSCILLATORS

Oscillators convert dc energy into rf. They differ from amplifiers in that the presence of an input signal is not essential for their operation while their output frequency and amplitude are primarily determined by their circuit characteristics. On the other hand, an amplifier produces its output signal only when the input signal is present.

Oscillators are inherently nonlinear and the simplest model has to take this effect into account; otherwise, the model will not reach a stable operating point. Keeping this in mind, we shall first construct a simple oscillator model and study its behavior, including the synchronization by a small signal injection. Then, we shall discuss the noise of free-running as well as synchronized oscillators. Since the concept of generalized functions and their Fourier analysis will be used extensively in the noise discussion, it is recommended for those who are not familiar with generalized functions that they read Appendix II before beginning Section 9.3.

9.1 A Simplified Oscillator Model

Let us consider a solid state oscillator which consists of a cavity containing an active device and a waveguide connecting the cavity and the load, as shown in Fig. 9.1. Since the active part of the device is generally very small compared to a wavelength in free space, the conventional voltage and current will be well defined at the terminals of the active part at the fundamental frequency of oscillation.

At each harmonic frequency, the device sees a certain environment, say

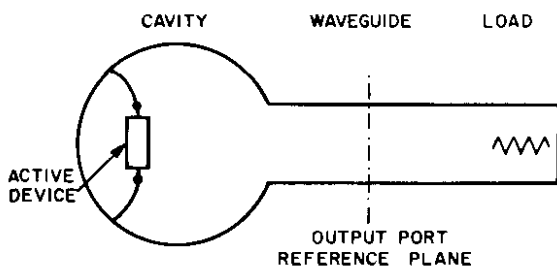


Fig. 9.1. A solid state oscillator.

impedance, which is not affected by the load condition if the harmonic content in the waveguide at the output port is small. Under this condition, the active part of the device at the fundamental frequency exhibits a certain nonlinear admittance which, in general, varies slowly with frequency (although this may not be true if anyone of the harmonics is in a strong resonance). Let us consider a two-port resonant cavity consisting of the region between an appropriate reference plane in the output waveguide of the oscillator and the terminals of the active part of the device. Suppose that only one resonant mode predominates, then the equivalent circuit of the cavity must be similar to that shown in Fig. 4.10. A certain length of transmission line may have to be inserted in the equivalent circuit before connecting it to the nonlinear device admittance since a shift of reference plane is not allowed on the device side of the equivalent circuit. This means that the nonlinear device admittance undergoes a transformation by the same transmission line before it is connected to the equivalent circuit in Fig. 4.10. Let us express the transformed nonlinear admittance in the form of a nonlinear impedance. Notice that the nonlinear impedance also varies slowly with frequency. The equivalent circuit of the oscillator, including the load resistance R_0 , then becomes a series connection of an inductance, a capacitance, a small positive resistance R_i representing the cavity loss, the nonlinear device impedance $-\bar{R} + j\bar{X}$, and R_0 as shown in Fig. 9.2. A small voltage source $e(t)$ represents the noise in the device and the circuit and/or the effect of a synchronizing signal, which we shall study later. The equation governing the behavior of the oscillator is given by

$$L \frac{di}{dt} + (R_i + R_0) i + \frac{1}{C} \int i dt + v = e(t) \quad (9.1)$$

where v is the voltage developed across the device impedance $-\bar{R} + j\bar{X}$.

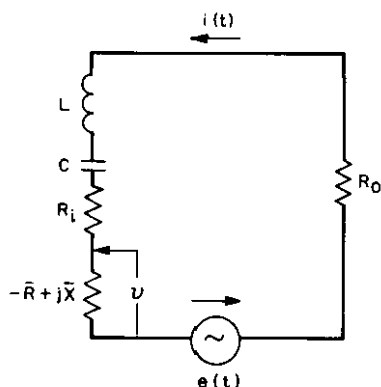


Fig. 9.2. Equivalent circuit of an oscillator.

When $e(t)$ is equal to zero, the current is expected to be approximately sinusoidal;

$$i(t) = A \cos(\omega t + \varphi) \quad (9.2)$$

where harmonic currents are neglected. The voltage drop across the device impedance can be expressed as

$$v = -\bar{R}A \cos(\omega t + \varphi) - \bar{X}A \sin(\omega t + \varphi) \quad (9.3)$$

provided that the frequency dependences of $\bar{R}$ and $\bar{X}$ are negligible. (A linear part of the device reactance can be included in the LC circuit of the cavity, thereby relaxing this restriction to a certain extent.) Now suppose that a small perturbation is given to the system by applying a small $e(t)$, then $i(t)$ may no longer be sinusoidal. In many cases, however, $i(t)$ is similar to the original waveform. We shall restrict ourselves to these cases only, and assume that $i(t)$ can be represented by

$$i(t) = A(t) \cos \{\omega t + \varphi(t)\} \quad (9.4)$$

where $A(t)$ and $\varphi(t)$ do not change appreciably over one cycle of the oscillation ($= 2\pi/\omega$). Under this condition, di/dt and $\int i dt$ are given by

$$\frac{di}{dt} = -A \left(\omega + \frac{d\varphi}{dt} \right) \sin(\omega t + \varphi) + \frac{dA}{dt} \cos(\omega t + \varphi) \quad (9.5)$$

$$\int i dt = \left(\frac{A}{\omega} - \frac{A}{\omega^2} \frac{d\varphi}{dt} \right) \sin(\omega t + \varphi) + \frac{1}{\omega^2} \frac{dA}{dt} \cos(\omega t + \varphi) \quad (9.6)$$

where use is made of integration by parts neglecting the higher order terms

of $(1/\omega)$. Equation (9.3) is still valid provided that A and φ are now interpreted as functions of time. Substituting Equations (9.3) through (9.6) into (9.1), multiplying by $\cos(\omega t + \varphi)$ or $\sin(\omega t + \varphi)$, and integrating with respect to time from $t - T_0$ to t , where T_0 is a period of the oscillation ($= 2\pi/\omega$), we obtain approximate differential equations in A and φ :

$$\left(L + \frac{1}{\omega^2 C}\right) \frac{dA}{dt} + (R_i + R_0 - \bar{R}) A = \frac{2}{T_0} \int_{t-T_0}^t e(t) \cos(\omega t + \varphi) dt \quad (9.7)$$

$$\left(-\omega L + \frac{1}{\omega C} - \bar{X}\right) - \left(L + \frac{1}{\omega^2 C}\right) \frac{d\varphi}{dt} = \frac{2}{AT_0} \int_{t-T_0}^t e(t) \sin(\omega t + \varphi) dt \quad (9.8)$$

where the orthogonality relation between sine and cosine functions has been used.

Even though higher harmonics at the output port are assumed to be small, since $e(t)$ is also small, their effect may be of the same order of magnitude in the original Eq. (9.1). However, the above integrations with respect to time, eliminate from our discussion the possible effect of higher harmonics because of the orthogonality relations between the fundamental and harmonic components.

Having obtained the differential equations which govern the behavior of A and φ , let us examine the free-running oscillation in these terms. For steady state free-running oscillation, $dA/dt = 0$ and $e(t) = 0$. Consequently, (9.7) gives

$$R_i + R_0 = \bar{R} \quad (9.9)$$

In general, $\bar{R}$ is a function of A and may appear similar to Fig. 9.3(a). From (9.9), the amplitude of the current is determined as the intersection of $\bar{R}$ and $R_i + R_0$. Let the amplitude be A_0 . If sR_0 represents A_0 times $-\partial\bar{R}/\partial A$ at A_0 , as shown in Fig. 9.3(a), then, for a small variation ΔA from A_0 , we have

$$R_i + R_0 - \bar{R} = sR_0 (\Delta A/A_0) \quad (9.10)$$

In the vicinity of A_0 , Eq. (9.7) becomes

$$\left(L + \frac{1}{\omega^2 C}\right) \frac{d\Delta A}{dt} + sR_0 \Delta A = 0$$

Suppose ΔA is not equal to zero at $t = 0$. If $s > 0$, then $|\Delta A|$ decreases exponentially with time, and A_0 represents a stable operating point. On the

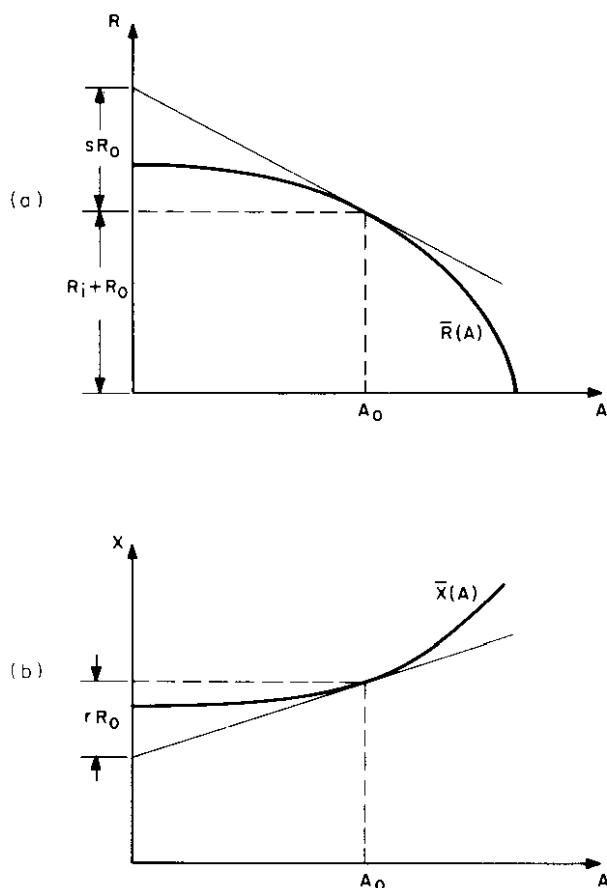


Fig. 9.3. (a) $\bar{R}$ as a function of A ; (b) $\bar{X}$ as a function of A .

other hand, if $\bar{R}$ intersects $R_i + R_0$ such that $s < 0$, then $|A|$ grows indefinitely with time. This means that such a point does not provide stable operation.

The steady state oscillation frequency is determined from (9.8) together with $d\phi/dt = 0$ and $e(t) = 0$;

$$\omega L + \bar{X} = (1/\omega C)$$

Assuming that $\bar{X}$ is small, the oscillation frequency is calculated to be

$$\omega_0' = \omega_0 \{1 - \frac{1}{2} Q_{ext}^{-1} (\bar{X}/R_0)\} \quad (9.11)$$

where ω_0 and Q_{ext} are given by

$$\omega_0 = (\sqrt{LC})^{-1/2}, \quad Q_{ext} = (\omega_0 L/R_0) \quad (9.12)$$

When $\bar{X}$ depends on A , as illustrated in Fig. 9.3(b), the value of $\bar{X}$ at $A = A_0$ must be used in (9.11) to obtain ω_0' .

In the above discussion, a small parallel conductance G has been neglected which represents the losses of other resonant modes and remains uncanceled after a proper shift of the reference plane at the output port. A similar conductance on the device side of the equivalent circuit can be considered to be included in the device impedance.

A plot of constant power and constant frequency contours on a Smith chart of the load impedance of an oscillator is called a Rieke diagram. When the load impedance has a reactive part X_0 , in (9.11) $\bar{X}$ has to be

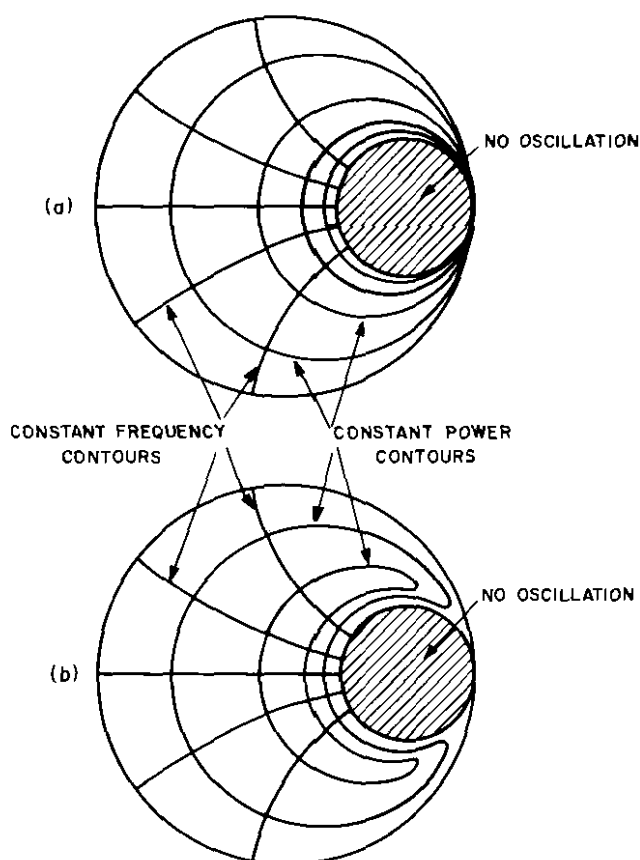


Fig. 9.4. Rieke diagram of an ideal oscillator (a) $G = 0$; (b) $G \neq 0$. The usual Smith chart constant resistance and reactance contours have been omitted.

replaced by $\bar{X} + X_0$. From this modified equation and (9.9), the Rieke diagram of our oscillator model has the appearance that is shown in Fig. 9.4(a). As the magnitude of the load impedance increases, the effect of the conductance G mentioned above is no longer negligible, and the constant power contours with nonzero G will be modified, as illustrated in Fig. 9.4(b). Although the constant power and constant frequency contours intersect each other perpendicularly in Fig. 9.4(a), their intersections will make oblique angles if $\bar{X}$ is nonlinear (Problem 9.2). Furthermore, if more than one resonant mode is strongly excited in the cavity, the Rieke diagram will not resemble the one in Fig. 9.4(a). How well our model represents a given oscillator in the vicinity of the operating point can easily be checked by measuring the power and frequency as functions of load impedance and drawing an appropriate Rieke diagram for the oscillator.

9.2 Synchronization by Injection Locking

In this section, we shall study the injection locking mechanism of oscillators. For simplicity, let us assume that R_0 is equal to the waveguide characteristic impedance Z_0 of the output port. Furthermore, let $e(t)$ in (9.7) and (9.8) represent the injection signal voltage $a_0 \cos \omega_s t$, where a_0 is assumed to be small. When the oscillator is synchronized, the oscillating frequency must be the same as ω_s . The right-hand side of (9.7) is given by

$$2T_0^{-1} \int_{t-\tau_0}^t a_0 \cos \omega_s t \cos(\omega_s t + \varphi) dt = a_0 \cos \varphi$$

Thus, (9.7) becomes

$$2L \frac{d \Delta A}{dt} + sR_0 \Delta A = a_0 \cos \varphi \quad (9.13)$$

where $\Delta A = A - A_0$, and use is made of $L + (1/\omega_s^2 C) \simeq L + (1/\omega_0^2 C) = 2L$. When the oscillation is in the steady state, $d \Delta A / dt = 0$, and (9.13) gives

$$\Delta A_0 = (a_0 \cos \varphi_0 / sR_0) \quad (9.14)$$

On the other hand, the right-hand side of (9.8) is given by

$$2(AT_0)^{-1} \int_{t-\tau_0}^t a_0 \cos \omega_s t \sin(\omega_s t + \varphi) dt = (a_0/A) \sin \varphi$$

Strictly speaking, A may not be the same as A_0 ; however, since a_0 is small,

the difference between A and A_0 can be neglected. Thus, (9.8) becomes

$$\left(-\omega_s L + \frac{1}{\omega_s C} - \bar{X} \right) - 2L \frac{d\varphi}{dt} = \frac{a_0}{A_0} \sin \varphi \quad (9.15)$$

For the steady state, $d\varphi/dt = 0$ and φ must be a constant. Let this constant be φ_0 and let ω_s be $\omega_0' + \Delta\omega_0$, where $\Delta\omega_0$ is the difference between the synchronizing and free-running frequencies. To a first order approximation, (9.15) now becomes

$$-\Delta\omega_0 2L = (a_0/A_0) \sin \varphi_0 + rR_0 (\Delta A_0/A_0)$$

where rR_0 is A_0 times $\partial\bar{X}/\partial A$ at A_0 , as shown in Fig. 9.3(b). Substituting (9.14) we have

$$\Delta\omega_0 = -\frac{1}{2}L^{-1} (a_0/A_0) \{1 + (r/s)^2\}^{1/2} \sin(\varphi_0 + \theta) \quad (9.16)$$

where

$$\theta = \tan(r/s) \quad (9.17)$$

Since $\sin(\varphi_0 + \theta)$ must be less than 1, the maximum value of $|\Delta\omega_0|$ for synchronization is given by

$$|\Delta\omega_0|_{\max} = \frac{1}{2}L^{-1} (a_0/A_0) \{1 + (r/s)^2\}^{1/2} \quad (9.18)$$

Using this value of $|\Delta\omega_0|_{\max}$, Eq. (9.16) becomes

$$\sin(\varphi_0 + \theta) = -(\Delta\omega_0/|\Delta\omega_0|_{\max}) \quad (9.19)$$

In the vicinity of A_0 and φ_0 , writing $A = A_0 + \Delta A_0 + \Delta A$ and $\varphi = \varphi_0 + \Delta\varphi$, Eqs. (9.13) and (9.15) reduce to

$$2L \frac{d\Delta A}{dt} + sR_0 \Delta A = -a_0 \Delta\varphi \sin \varphi_0 \quad (9.20)$$

$$-2L \frac{d\Delta\varphi}{dt} - rR_0 \frac{\Delta A}{A_0} = \frac{a_0}{A_0} \Delta\varphi \cos \varphi_0 \quad (9.21)$$

Eliminating ΔA from these equations, we obtain

$$\begin{aligned} \frac{d^2 \Delta\varphi}{dt^2} + \frac{1}{2L} \left(sR_0 + \frac{a_0}{A_0} \cos \varphi_0 \right) \frac{d\Delta\varphi}{dt} \\ + \frac{sR_0}{4L^2} \frac{a_0}{A_0} \left\{ 1 + \left(\frac{r}{s} \right)^2 \right\}^{1/2} \cos(\varphi_0 + \theta) \Delta\varphi = 0 \end{aligned} \quad (9.22)$$

For stable operation, $|\Delta\varphi|$ must decrease with time if $\Delta\varphi$ is not equal to

zero at $t = 0$. Thus, we have two conditions

$$sR_0 + (a_0/A_0) \cos \varphi_0 > 0, \quad sR_0 (a_0/A_0) \{1 + (r/s)^2\}^{1/2} \cos(\varphi_0 + \theta) > 0$$

The first condition is usually satisfied when a_0 is small and sR_0 is positive as required for stable free-running oscillation. The second condition is then equivalent to

$$\cos(\varphi_0 + \theta) > 0 \quad (9.23)$$

For a given $\Delta\omega_0 = \omega_s - \omega_0'$, φ_0 is uniquely determined from (9.19) and (9.23).

The above discussions were in terms of voltage and current. In the micro-wave circuit, power is a more meaningful quantity; therefore, let us express the synchronizing range in terms of power. The available power from the synchronizing source is given by

$$P_s = (a_0/\sqrt{2})^2/4R_0 = a_0^2/8R_0$$

while the output power of the free-running oscillator is given by

$$P_0 = R_0 (A_0/\sqrt{2})^2 = \frac{1}{2} R_0 A_0^2$$

Substituting these into (9.18), we have

$$|\Delta\omega_0|_{\max} = \frac{\omega_0}{Q_{\text{ext}}} \left\{ 1 + \left(\frac{r}{s} \right)^2 \right\}^{1/2} \left(\frac{P_s}{P_0} \right)^{1/2} \quad (9.24)$$

where the external Q is defined by (9.12).

The injection of a synchronizing signal and the extraction of output power from an oscillator are generally done through a circulator matched to the waveguide characteristic impedance. In this case, the output power is not given by $R_0 A_0^2/2$ but is given by the power emerging from the oscillator through the circulator, i.e., the actual power to R_0 . From the discussion in Section 1.4, the output power becomes

$$P = \frac{1}{4} R_0^{-1} |V - R_0 I|^2$$

where V is the voltage across the terminals of R_0 (including $e(t)$) given by

$$V = E - R_0 I = (a_0/\sqrt{2}) - R_0 \{(A_0 + \Delta A_0)/\sqrt{2}\} e^{j\varphi_0}$$

Thus, to a first order approximation, we have

$$P = P_0 + A_0 a_0 \cos \varphi_0 \left(\frac{1}{s} - \frac{1}{2} \right) = P_0 \left\{ 1 + 2 \left(\frac{P_s}{P_0} \right)^{1/2} \left(\frac{2}{s} - 1 \right) \cos \varphi_0 \right\} \quad (9.25)$$

When combined with (9.19) this equation tells us how P changes with $\Delta\omega_0 = \omega_s - \omega_0'$. If the circuit is adjusted so that the maximum output is obtained under the free-running condition, one of the constant output power contours ($\frac{1}{2}R_0A^2 = \text{const}$) plotted on the plane of Fig. 9.3(a) should be tangent to the curve $\bar{R}$ at A_0 . Since the constant output power contours are represented by

$$\frac{1}{2}R_0A^2 = \frac{1}{2}(\bar{R} - R_i)A^2 = \text{const}$$

the maximum power condition is given by

$$(\partial\bar{R}/\partial A)A^2 + 2A(\bar{R} - R_i) = 0$$

or equivalently

$$(\partial\bar{R}/\partial A) = -(2R_0/A)$$

Remembering that s is defined as $-A_0/R_0$ times $\partial\bar{R}/\partial A$ at the operating point, the above result shows that when the circuit is adjusted for maximum power, s is equal to two, and that the variation in the output power of our model due to the injection of synchronizing signal becomes zero to a first order approximation.

9.3 Noise in Oscillators

For the discussion of noise, it is necessary to evaluate the integrals on the right-hand sides of (9.7) and (9.8) for the noise voltage $e(t)$. As a simple noise model, let us assume that $e(t)$ consists of a large number of elementary pulses each of which is represented by $\varepsilon\delta(t - t_0)$, where ε gives the strength of the pulse and t_0 the time of its occurrence. Both ε and t_0 are independent random variables from one pulse to the next. Using (II.20) in Appendix II, the autocorrelation function of $e(t)$ is given by

$$R_e(\tau) = n \langle \varepsilon^2 \rangle \delta(\tau) \quad (9.26)$$

where n is the average number of pulses in a unit time, and $\langle \varepsilon^2 \rangle$ is the ensemble average of the square of the pulse strength ε . From (9.26), the power spectral density of $e(t)$ is calculated to be

$$|e(f)|^2 = n \langle \varepsilon^2 \rangle \equiv |e|^2 \quad (9.27)$$

which is independent of frequency. When the spectral density is independent of frequency, the noise is said to be white. The thermal noise in resistors and

the shot-noise in vacuum tubes and p - n junctions are both considered to be white for most practical applications.

Since Eqs. (9.7) and (9.8) are linear, the principle of superposition holds. Let us first consider the effect of just one pulse represented by $\epsilon\delta(t-t_0)$; the integral

$$2T_0^{-1} \int_{t-T_0}^t e(t) \cos(\omega t + \phi) dt \quad (9.28)$$

becomes a rectangular pulse with constant height $(2\epsilon/T_0) \cos(\omega t_0 + \phi)$ and the length T_0 as shown in Fig. 9.5. Since T_0 is considered short for any

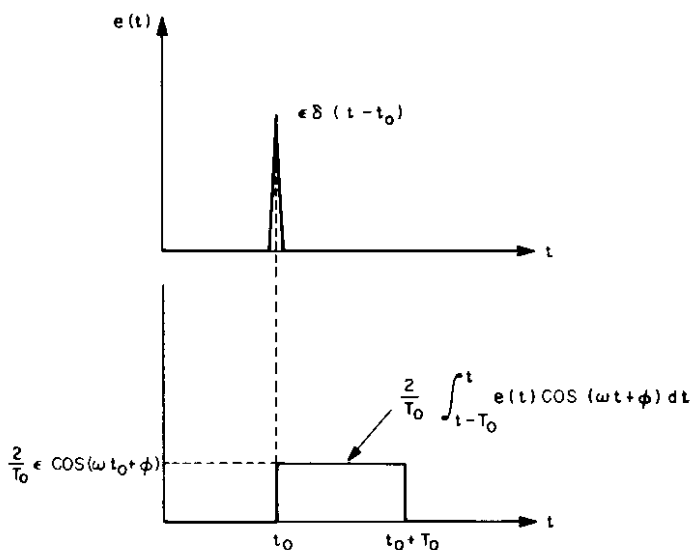


Fig. 9.5. The relation between $2T_0^{-1} \int_{t-T_0}^t e(t) \cos(\omega t + \phi) dt$ and $e(t)$ for a single pulse in $e(t)$.

variation in A , this pulse can be considered as another pulse having the form of a δ -function with strength $2\epsilon \cos(\omega t_0 + \phi)$ located at $t = t_0$ for the practical purpose of calculating the effect using (9.7). If $n_1(t)$ indicates the total effect of the integral (9.28), $n_1(t)$ consists of many such pulses. Since the autocorrelation function of $n_1(t)$ is easily calculated to be $2n\langle\epsilon^2\rangle\delta(\tau)$, we have

$$|n_1(f)|^2 = 2|e|^2 \quad (9.29)$$

Similarly, indicating the effect of

$$2T_0^{-1} \int_{t-T_0}^t e(t) \sin(\omega t + \varphi) dt$$

by $n_2(t)$, we obtain

$$|n_2(f)|^2 = 2|e|^2 \quad (9.30)$$

Furthermore, because of the orthogonality between sine and cosine functions, the cross-correlation between $n_1(t)$ and $n_2(t)$ is calculated to be zero.

We are now in a position to investigate noise in free-running oscillators. In the following discussion it will be assumed that $r = 0$ in order to simplify some calculations which otherwise become very lengthy.

For the amplitude fluctuation we write $\Delta A = A - A_0$, and from (9.7)

$$2L(d\Delta A/dt) + sR_0 \Delta A = n_1(t)$$

Calculating the power spectral densities of both sides and dividing by $4\omega^2 L^2 + s^2 R_0^2$ we obtain

$$|\Delta A(f)|^2 = \frac{|n_1(f)|^2}{4\omega^2 L^2 + s^2 R_0^2} = \frac{2|e|^2}{4\omega^2 L^2 + s^2 R_0^2} \quad (9.31)$$

If the output of the oscillator is detected by an envelope detector, we expect to observe the output spectrum given by (9.31). Strictly speaking, part of $e(t)$ comes from the noise source in R_0 which introduces a slight modification in the discussion of the output power measurement as was done in Section 9.2. However, for most cases the difference may be negligible.

For the phase fluctuation, setting $r = 0$ and $\omega = \omega_0'$, (9.8) becomes

$$-2L(d\varphi/dt) = n_2(t)/A_0$$

which gives

$$|\varphi(f)|^2 = \frac{|n_2(f)|^2}{4\omega^2 L^2 A_0^2} = \frac{2|e|^2}{4\omega^2 L^2 A_0^2} \quad (9.32)$$

If the output of the oscillator is detected by an FM discriminator following an ideal limiter, the output should be proportional to $d\varphi/dt$ which has a white spectrum in this ideal case of $r = 0$.

An ordinary spectrum analyzer shows the spectral density of the output power directly. The best way to calculate it may be to calculate the autocorrelation function of the current waveform and transform the result to the spectral density by a Fourier transformation. Since $n_1(t)$ and $n_2(t)$ have no correlation, $A(t)$ and $\varphi(t)$ are uncorrelated, and the autocorrelation function of $R_i(\tau)$ becomes

$$R_i(\tau) = \langle A(t) A(t+\tau) \cos(\omega_0 t + \varphi(t)) \cos(\omega_0 t + \omega_0 \tau + \varphi(t+\tau)) \rangle \\ = \frac{1}{2} (A_0^2 + R_{AA}(\tau)) \cos \omega_0 \tau \langle \cos(\varphi(t+\tau) - \varphi(t)) \rangle \quad (9.33)$$

where $\langle \rangle$ indicates the ensemble average and use is made of $\langle \sin(\varphi(t+\tau) - \varphi(t)) \rangle = 0$.

Once the power spectral density of ΔA is available, the calculation of $R_{AA}(\tau)$ is straightforward. From (9.31), we have

$$R_{AA}(\tau) = \int_{-\infty}^{\infty} |\Delta A(f)|^2 e^{i\omega\tau} df = \frac{|e|^2}{2LsR_0} \exp\left(-\frac{sR_0}{2L} |\tau|\right) \quad (9.34)$$

Since the calculation of $\langle \cos(\varphi(t+\tau) - \varphi(t)) \rangle$ is involved, let us first consider $\langle (\varphi(t+\tau) - \varphi(t))^2 \rangle$. If $u(z)$ is unity from $z = -\tau$ to $z = 0$ and is zero elsewhere, then referring to Fig. 9.6, $x(t) = \varphi(t+\tau) - \varphi(t)$ can be expressed as

$$x(t) = \int_{-\infty}^{\infty} \frac{d\varphi(\theta)}{d\theta} u(t-\theta) d\theta = \int_{-\infty}^{\infty} \left(\frac{-1}{2LA_0} \right) n_2(\theta) u(t-\theta) d\theta \quad (9.35)$$

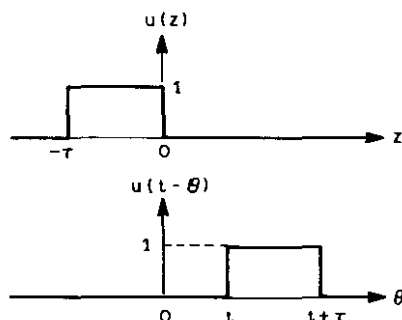


Fig. 9.6. The relation between $u(z)$ and $u(t - \theta)$.

The Fourier transform $X_T(f)$ of $x_T(t)$ (see Appendix II for the definition of $X_T(f)$ and $x_T(t)$) is therefore given by the product of the Fourier transforms of $n_{2T}(t)$ and $u(t)$ multiplied by $-1/2LA_0$. The power spectral density of $x(t)$ is given by

$$|x(f)|^2 = \lim_{T \rightarrow \infty} \frac{1}{2T} |X_T(f)|^2 = \frac{1}{4L^2 A_0^2} |n_2(f)|^2 |U(f)|^2 \\ = \frac{|e|^2}{2L^2 A_0^2} \frac{2(1 - \cos \omega\tau)}{\omega^2} \quad (9.36)$$

From (9.36), $R_x(0) = \langle (\varphi(t+\tau) - \varphi(t))^2 \rangle$ can be calculated as follows:

$$R_x(0) = \int_{-\infty}^{\infty} |x(f)|^2 df = \frac{|e|^2}{L^2 A_0^2} \int_{-\infty}^{\infty} \frac{1 - \cos \omega \tau}{\omega^2} df = \frac{|e|^2}{L^2 A_0^2} \frac{|\tau|}{2} \quad (9.37)$$

where we have used

$$\begin{aligned} \int_{-\infty}^{\infty} \frac{1 - \cos \omega \tau}{\omega^2} df &= \int_{-\infty}^{\infty} \int_0^{|\tau|} \frac{\sin \omega \theta}{\omega} d\theta df \\ &= \frac{1}{2} \int_{-\infty}^{\infty} \int_0^{|\tau|} \int_{-\theta}^{\theta} \cos \omega \xi d\xi d\theta df \\ &= \frac{1}{2} \int_0^{|\tau|} \int_{-\theta}^{\theta} \int_{-\infty}^{\infty} \cos \omega \xi df d\xi d\theta \\ &= \frac{1}{2} \int_0^{|\tau|} \int_{-\theta}^{\theta} \delta(\xi) d\xi d\theta = \frac{|\tau|}{2} \end{aligned}$$

Noting that $\varphi(t+\tau) - \varphi(t)$ has the Gaussian distribution and that $\sigma^2 = \langle (\varphi(t+\tau) - \varphi(t))^2 \rangle$, we can now calculate $\langle \cos(\varphi(t+\tau) - \varphi(t)) \rangle$ as follows:

$$\begin{aligned} \langle \cos(\varphi(t+\tau) - \varphi(t)) \rangle &= \int_{-\infty}^{\infty} (2\pi\sigma^2)^{-1/2} \exp\left(-\frac{\varphi^2}{2\sigma^2}\right) \cos \varphi d\varphi \\ &= \exp\left(-\frac{\sigma^2}{2}\right) = \exp\left(-\frac{|e|^2}{4L^2 A_0^2} |\tau|\right) \end{aligned} \quad (9.38)$$

Substituting (9.34) and (9.38) into (9.33), we obtain

$$R_i(\tau) = \frac{1}{2} \left\{ A_0^2 + \frac{|e|^2}{2LsR_0} \exp\left(-\frac{sR_0}{2L} |\tau|\right) \right\} \exp\left(-\frac{|e|^2}{4L^2 A_0^2} |\tau|\right) \cos \omega_0 \tau$$

The power spectral density of $i(t)$ can now be obtained by the Fourier transformation of the autocorrelation function. Noting that $sR_0/2L \gg |e|^2/4L^2 A_0^2$, we have

$$\begin{aligned} |i(f)|^2 &= \int_{-\infty}^{\infty} R_i(\tau) e^{-j\omega\tau} d\tau \\ &= \frac{|e|^2}{8L^2} \left[\left\{ (\omega - \omega_0)^2 + \left(\frac{|e|^2}{4L^2 A_0^2} \right)^2 \right\}^{-1} + \left\{ (\omega + \omega_0)^2 + \left(\frac{|e|^2}{4L^2 A_0^2} \right)^2 \right\}^{-1} \right] \\ &\quad + \frac{|e|^2}{8L^2} \left[\left\{ (\omega - \omega_0)^2 + \left(\frac{sR_0}{2L} \right)^2 \right\}^{-1} + \left\{ (\omega + \omega_0)^2 + \left(\frac{sR_0}{2L} \right)^2 \right\}^{-1} \right] \end{aligned}$$

Neglecting the contribution from the noise source in R_0 , $2R_0|i(f)|^2$ gives the output power spectral density function $P(f)$ ($f > 0$) for the oscillator. The factor 2 in front of $R_0|i(f)|^2$ comes from the fact that both the positive and negative parts of the frequency spectrum contribute to $P(f)$. Since $\omega - \omega_0 \ll \omega + \omega_0$, we have

$$P(f) = \frac{\omega_0^2}{2Q_{\text{ext}}^2} N \times \left\{ \frac{1}{(\omega - \omega_0)^2 + (\omega_0^2/4Q_{\text{ext}}^2)^2 (N/P_0)^2} + \frac{1}{(\omega - \omega_0)^2 + (\omega_0/Q_{\text{ext}})^2 (s/2)^2} \right\} \quad (9.39)$$

where N is given by

$$N = |e|^2/2(\bar{R} - R_i) \quad (9.40)$$

It is interesting to note that if an amplifier were built using the same negative resistance ($\bar{R} - R_i$) with the same noise voltage $e(t)$, the optimum noise measure obtainable would be given by

$$M_{\text{opt}} = (N/kT_i)$$

where k is the Boltzmann constant.

In (9.39), $P(f)$ consists of two terms, the first is called the FM noise since a spectrum analyzer essentially shows this spectrum if the output of the oscillator is displayed after eliminating the amplitude fluctuation by a limiter. The second term is called the AM noise. Since the relation

$$\frac{1}{2}\omega_0 Q_{\text{ext}}^{-1} (N/P_0) \ll s$$

usually holds, the first term predominates over the second when ω is close to ω_0 . For large $\omega - \omega_0$, however, both terms give equal contributions.

Let us next consider the noise in oscillators synchronized by injection locking. In the case of free-running oscillators, φ could take almost any value; when synchronized, however, φ should stay in the vicinity of φ_0 because of the restoring force due to the synchronizing signal. Writing $A = A_0 + \Delta A$, $\varphi = \varphi_0 + \Delta\varphi$, and assuming ΔA , ΔA_0 , and $\Delta\varphi$ are all small and $r = 0$, Eqs. (9.7) and (9.8) become

$$\begin{aligned} 2L(d\Delta A/dt) + sR_0\Delta A + a_0\Delta\varphi \sin\varphi_0 &= n_1(t) \\ -2L(d\Delta\varphi/dt) - (a_0/A_0)\Delta\varphi \cos\varphi_0 &= \{n_2(t)/A_0\} \end{aligned}$$

respectively. From these equations, we obtain

$$|\Delta A(f)|^2 = \frac{\{n_1(f)\}^2}{4\omega^2 L^2 + s^2 R_0^2} + \frac{a_0^2 \sin^2 \varphi_0}{4\omega^2 L^2 + s^2 R_0^2} \frac{\{n_2(f)\}^2}{4\omega^2 L^2 A_0^2 + a_0^2 \cos^2 \varphi_0} \\ = \frac{2|e|^2}{4\omega^2 L^2 + s^2 R_0^2} \frac{4\omega^2 L^2 A_0^2 + a_0^2}{4\omega^2 L^2 A_0^2 + a_0^2 \cos^2 \varphi_0} \quad (9.41)$$

$$|\Delta \varphi(f)|^2 = \frac{\{n_2(f)\}^2}{4\omega^2 L^2 A_0^2 + a_0^2 \cos^2 \varphi_0} = \frac{2|e|^2}{4\omega^2 L^2 A_0^2 + a_0^2 \cos^2 \varphi_0} \quad (9.42)$$

It can be seen by noting the presence of a_0 that the phase fluctuation is considerably reduced, but the envelope fluctuation becomes slightly larger with synchronization.

In order to calculate the output power spectral density function $P(f)$, let us first calculate $R_i(\tau)$. Noting that both ΔA and $\Delta \varphi$ are small, we have

$$R_i(\tau) = \frac{1}{2}(A_0 + \Delta A_0)^2 (1 - \langle \Delta \varphi^2 \rangle) \cos \omega_s \tau \\ + \frac{1}{2} A_0^2 R_{\Delta \varphi}(\tau) \cos \omega_s \tau + \frac{1}{2} R_{\Delta A}(\tau) \cos \omega_s \tau \\ + \frac{1}{2} A_0 \sin \omega_s \tau (\langle \Delta A(\tau) \Delta \varphi(0) \rangle - \langle \Delta A(0) \Delta \varphi(\tau) \rangle)$$

From this, we can obtain the power spectral density of $i(t)$ by a standard procedure. If we neglect the contribution from the noise source in R_0 and calculate the power emerging from the oscillator following the method described in the end of Section 9.2, we obtain

$$P(f) = \left\{ P \left(1 - \frac{N}{4(P_0 P_s)^{1/2}} \frac{\omega_0}{Q_{\text{ext}}} \frac{1}{\cos \varphi_0} \right) - \frac{N \omega_0}{4 Q_{\text{ext}}} \right\} \delta(f - f_s) \\ + \frac{\omega_0^2}{2 Q_{\text{ext}}^2} N \left\{ \frac{1}{(\omega - \omega_s)^2 + |\Delta \omega_0|_{\text{max}}^2 \cos^2 \varphi_0} \right. \\ + \frac{1}{(\omega - \omega_s)^2 + (\omega_0/Q_{\text{ext}})^2 (s/2)^2} \frac{(\omega - \omega_s)^2 + |\Delta \omega_0|_{\text{max}}^2}{(\omega - \omega_s)^2 + |\Delta \omega_0|_{\text{max}}^2 \cos^2 \varphi_0} \\ \left. + \frac{\omega - \omega_s}{(\omega - \omega_s)^2 + (\omega_0/Q_{\text{ext}})^2 (s/2)^2} \frac{2 |\Delta \omega_0|_{\text{max}} \sin \varphi_0}{(\omega - \omega_s)^2 + |\Delta \omega_0|_{\text{max}}^2 \cos^2 \varphi_0} \right\} \quad (9.43)$$

where P is given by (9.25). The first term is a pure spectrum. The second term represents the noise component which in turn consists of three terms. The first term gives the FM noise, the second term corresponds to the AM noise, and the last term represents the interaction between the FM and AM noise. As can be seen by comparison of (9.39) and (9.43), the FM noise is

of the same order of magnitude as in the free-running case when $(\omega - \omega_s)^2$ is large. However, as ω approaches ω_s , a considerable improvement is obtained by synchronization compared to the free-running case (9.39). On the other hand, the AM noise is not improved by synchronization. As a matter of fact, since $\cos^2 \varphi_0$ is smaller than unity, the AM noise increases slightly. If $\cos^2 \varphi_0$ becomes small, i.e., if the synchronizing frequency is near the edge of the synchronizing range, both the FM and the AM noise become extremely large. In the limit of $\cos^2 \varphi_0$ approaching zero, however, the assumption of ΔA and $\Delta \varphi$ being small is no longer valid, and the above result should not be applied.

The interaction term shifts the noise from one side of ω_s to the other if ω_s and ω_0 do not coincide, and as a result, the noise spectrum becomes unsymmetrical with respect to ω_s . The sign of the interaction term is such that the ω_0 side of ω_s always becomes noisier.

So far the contribution from $1/f$ noise has been completely neglected. In practice, the circuit parameters of the active device are expected to fluctuate slowly with time, and the spectra of the fluctuations will be nearly proportional to $1/f$, according to many experiments. This type of noise does not play a significant role in the case of ordinary high frequency amplifiers. However, in the case of oscillators, the situation is different. As shown in Section 9.1., the behavior of the amplitude and phase of an oscillating current is governed by (9.7) and (9.8). In these equations, $\bar{R}$ and $\bar{X}$ are expected to fluctuate with time. Concentrating on (9.7), if $\bar{R}$ fluctuates due to the fluctuation of the device resistance, then transfer the fluctuating part $\Delta \bar{R}$ to the right-hand side. If we write $\Delta \bar{R}(t) = n_R(t)/A$, $n_R(t)$ has the dimension of voltage, and the equation becomes

$$\left(L + \frac{1}{\omega^2 C}\right) \frac{dA}{dt} + (R_i + R_0 - \bar{R}) A = \frac{2}{T_0} \int_{t-T_0}^t e(t) \cos(\omega t + \varphi) dt + n_R(t) \quad (9.44)$$

Similarly, (9.8) becomes

$$\begin{aligned} & \left(-\omega L + \frac{1}{\omega C} - \bar{X}\right) - \left(L + \frac{1}{\omega^2 C}\right) \frac{d\varphi}{dt} \\ & = \frac{2}{AT_0} \int_{t-T_0}^t e(t) \sin(\omega t + \varphi) dt + \frac{n_x(t)}{A} \end{aligned} \quad (9.45)$$

where $n_x(t) = A \Delta \bar{X}(t)$.

Equations (9.44) and (9.45) are the fundamental equations for A and φ

when $1/f$ noise is taken into account. The terms $|n_R(f)|^2$ and $|n_x(f)|^2$ are both proportional to $1/f$. In general, $n_R(t)$ and $n_x(t)$ are independent of $n_1(t)$ and $n_2(t)$, the white noise we previously introduced. However, $n_R(t)$ and $n_x(t)$ may be cross-correlated with each other, and this makes the calculation of $|i(f)|^2$ complicated even when $r=0$. On the other hand, the spectrum calculation of phase and amplitude fluctuations remains straightforward. For instance, in the case of free-running oscillators the results for $|AA(f)|^2$ and $|\varphi(f)|^2$ remain valid if $2|e|^2$ is simply replaced by $2|e|^2 + |n_R(f)|^2$ and $2|e|^2 + |n_x(f)|^2$, respectively. It is worth noting that the spectrum of frequency fluctuation $|d\varphi(f)/dt|^2$ is no longer white. Depending on the relative magnitudes of $2|e|^2$ and $|n_x(f)|^2$, the spectrum exhibits a certain slope, with the steepest region being about 3 dB per octave.

The calculated result of the phase fluctuation for synchronized oscillators also remains valid if $2|e|^2$ is replaced by $2|e|^2 + |n_x(f)|^2$. Phase fluctuations can therefore be greatly reduced by synchronization, even if $1/f$ noise is taken into account. The amplitude fluctuation contains a new interaction term between $n_R(t)$ and $n_x(t)$, and whether or not this fluctuation is reduced by injection locking depends on the correlation between the two. If they are uncorrelated, the amplitude fluctuation will be increased by injection locking as before.

PROBLEMS

- 9.1 Assuming that the reactive part of the device impedance is a function of the current amplitude, draw a Rieke diagram corresponding to Fig. 9.4(a).
- 9.2 Calculate the angle with which the constant power and constant frequency contours intersect each other at the operating point in the previous problem.
- 9.3 Prove that generalized functions corresponding to two different continuous functions are not equal. Use $\psi_{\alpha\beta}(z)$ defined by (II.16) in Appendix II to test the equality.
- 9.4 Calculate the autocorrelation function of $A \cos(\omega t + \varphi)$ and its power spectral density.
- 9.5 Show that the power spectral density of shot noise current is given by

$$|i(f)|^2 = eI_0$$

where I_0 is the average current and e is the magnitude of electronic charge. Note that this gives (7.71) when the contribution from the negative frequency part of the spectrum is taken into account; i.e., $\langle i^2 \rangle = 2|i(f)|^2 B$.

- 9.6 Suppose that the synchronizing signal is noisy and represented by $(a_0 + \Delta a) \cos(\omega_s t + \psi)$, where Δa and ψ indicate the amplitude and phase fluctuation.

tuations, respectively. Show that the power spectral densities of the amplitude and phase fluctuations of the synchronized oscillator are given by

$$\begin{aligned}
 |\Delta A(f)|^2 &= R_0^{-1} \{(\omega/\omega_0)^2 Q_{\text{ext}}^2 + (s/2)^2\}^{-1} \{\omega^2 + |\Delta\omega_0|_{\text{max}}^2 \cos^2 \varphi_0\}^{-1} \\
 &\quad \times \{|\Delta\omega_0|_{\text{max}}^2 \cos^2 \varphi_0 2P_S \sin^2 \varphi_0 |\psi(f)|^2 + (\omega^2 + |\Delta\omega_0|_{\text{max}}^2) N\} \\
 |(\varphi - \varphi_0)(f)|^2 &= \{\omega^2 + |\Delta\omega_0|_{\text{max}}^2 \cos^2 \varphi_0\}^{-1} \left\{ |\Delta\omega_0|_{\text{max}}^2 \cos^2 \varphi_0 |\psi(f)|^2 + \frac{\omega_0^2}{2Q_{\text{ext}}^2} \frac{N}{P_0} \right\}
 \end{aligned}$$

where it is assumed that $P_0 \gg P_S$, $a_0 \gg \Delta a$, and $r = 0$.

9.7 Show that $|\Delta A(f)|^2$ and $|\varphi(f)|^2$ of a free-running oscillator become

$$\begin{aligned}
 |\Delta A(f)|^2 &= NR_0^{-1} \{(\omega/\omega_0)^2 Q_{\text{ext}}^2 + (s/2)^2\}^{-1} \\
 |\varphi(f)|^2 &= (N/P_0) \{2Q_{\text{ext}}^2 (\omega/\omega_0)^2\}^{-1} [1 + (r/2)^2 \{(\omega/\omega_0)^2 Q_{\text{ext}}^2 + (s/2)^2\}^{-1}]
 \end{aligned}$$

when the nonlinearity of the device reactance is taken into account, i.e., $r \neq 0$.



APPENDIX I

PROOF OF $\lim_{n \rightarrow \infty} k_n^2 = \infty$

In this appendix we shall prove the infinite growth of the eigenvalues which we used in the discussion concerning the completeness of eigenfunctions.

1. One-Dimensional Case

The proof will be presented in several steps.

LEMMA 1. Let f and g be single-valued real functions of x , then

$$\int f^2 dx \int g^2 dx \geq \left(\int fg dx \right)^2 \quad (\text{I.1})$$

where the integrals are from $x = a$ to b ($b > a$).

Proof. Noting that the square of a real function is positive, we have

$$\begin{aligned} 0 &\leq \int \left(f \int g^2 dx - g \int fg dx \right)^2 dx \\ &= \int f^2 dx \left(\int g^2 dx \right)^2 - 2 \left(\int fg dx \right)^2 \int g^2 dx + \int g^2 dx \left(\int fg dx \right)^2 \\ &= \int g^2 dx \left\{ \int f^2 dx \int g^2 dx - \left(\int fg dx \right)^2 \right\} \end{aligned}$$

Since the integral of g^2 is positive, the terms inside the brackets on the right-hand side must be positive and the proof of (I.1) is complete.

LEMMA 2. Let f and g be single-valued real functions of x , then

$$\int (f \pm g)^2 dx \leq 2 \left(\int f^2 dx + \int g^2 dx \right) \quad (\text{I.2})$$

where the integrals are from $x = a$ to b ($b > a$)

Proof. From

$$0 \leq \int (f \pm g)^2 dx = \int f^2 dx \pm 2 \int fg dx + \int g^2 dx$$

we have

$$\left| 2 \int fg dx \right| \leq \int f^2 dx + \int g^2 dx \quad (\text{I.3})$$

Substituting (I.3) into

$$\int (f \pm g)^2 dx \leq \int f^2 dx + 2 \left| \int fg dx \right| + \int g^2 dx$$

we obtain (I.2).

LEMMA 3. Let f be a single-valued function of x whose derivative is square-integrable over L .

Let us divide L into N sections and indicate the i th section by L_i . Let $\bar{f}(i)$ be the average of f over L_i , then

$$\bar{f}(i) = L_i^{-1} \int_{L_i} f dx$$

A staircase function can be defined over L such that its value in each L_i is equal to $\bar{f}(i)$. If this staircase function is indicated by $\bar{f}$, then

$$\int_L (f - \bar{f})^2 dx \leq K \int_L (df/dx)^2 dx \quad (\text{I.4})$$

where K is a constant which approaches zero as $N \rightarrow \infty$ and $\max L_i \rightarrow 0$.

Proof. Since f is continuous, its value becomes equal to $\bar{f}(i)$ somewhere in L_i . Let x_0 be such a point, then we have

$$f - \bar{f}(i) = \int_{x_0}^x (df/dx) dx$$

in L_i . Squaring both sides and applying (I.1) with f being replaced by

df/dx , and g by 1, we obtain

$$\begin{aligned} \{f - \bar{f}(i)\}^2 &= \left\{ \int_{x_0}^x (df/dx) dx \right\}^2 \\ &\leq \left| \int_{x_0}^x dx \int_{x_0}^x (df/dx)^2 dx \right| \leq L_i \int_{L_i} (df/dx)^2 dx \end{aligned}$$

It follows from this that

$$\int_L (f - \bar{f})^2 dx \leq L \max L_i \int_L (df/dx)^2 dx$$

If $\max L_i \rightarrow 0$ as $N \rightarrow \infty$, then $L \max L_i$ approaches zero. Writing $L \max L_i$ as K , (I.4) is proved.

LEMMA 4. An infinite sequence $a_1, a_2, \dots, a_n, \dots$ satisfying

$$|a_n| \leq C \quad (n = 1, 2, \dots)$$

contains at least one subsequence which converges to a finite constant. The constant C in the above inequality is fixed and independent of n .

Proof. Bisect the section from $-C$ to $+C$ at its center and classify the a_n 's into two groups, one satisfying $-C \leq a_n < 0$, and the other satisfying $0 \leq a_n \leq C$. At least one of the two groups should contain an infinite number of a_n 's. Again let us bisect the section containing this group at its center, then at least one of these two sections contains an infinite number of a_n 's, as before. Continuing this process indefinitely, we see that there is a constant α somewhere between $-C$ and $+C$ such that $\alpha \pm \varepsilon$ contains an infinite number of a_n 's no matter how small ε is. Let $a(m)$ be the first a_n to appear within the range $\alpha \pm (1/m)$ where m is a positive integer. By increasing m sequentially, an infinite sequence $\dots, a(m), a(m+1), \dots$ will be obtained. If a particular $a(m)$ happens to be the same a_n as previously used for $a(m-1)$, this $a(m)$ will be omitted. The sequence thus obtained is a subsequence of the a_n 's, and it converges to α . This completes the proof.

RELICH'S THEOREM. A sequence of functions f_n ($n = 1, 2, \dots, \infty$), each satisfying

$$\int_L f_n^2 dx \leq C, \quad \int_L (df_n/dx)^2 dx \leq C \quad (\text{I.5})$$

contains a subsequence such that

$$\lim_{m, n \rightarrow \infty} \int_L (f_n - f_m)^2 dx = 0$$

Proof. The proof will be carried out in two steps. Divide L into N sections as we did in Lemma 3, and let $\bar{f}$ be the staircase function whose value within L_i is equal to $\bar{f}(i)$. In the first step, we shall prove that

$$\lim_{N \rightarrow \infty} \left| \int \{(f_n - f_m)^2 - (\bar{f}_n - \bar{f}_m)^2\} dx \right| \rightarrow 0 \quad (I.6)$$

where $\max L_i \rightarrow 0$ is assumed as $N \rightarrow \infty$.

For simplicity, we write

$$\varphi = f_n - f_m, \quad \bar{\varphi} = \bar{f}_n - \bar{f}_m$$

then we obtain

$$\begin{aligned} \left| \int (\varphi^2 - \bar{\varphi}^2) dx \right| &\leq \int |\varphi^2 - \bar{\varphi}^2| dx \leq \int |\varphi + \bar{\varphi}| |\varphi - \bar{\varphi}| dx \\ &\leq \left\{ \int (\varphi + \bar{\varphi})^2 dx \int (\varphi - \bar{\varphi})^2 dx \right\}^{1/2} \end{aligned} \quad (I.7)$$

where use is made of (I.1) for the last inequality. For the first term in the square root, Lemma 2 gives

$$\int (\varphi + \bar{\varphi})^2 dx \leq 2 \left(\int \varphi^2 dx + \int \bar{\varphi}^2 dx \right) \quad (I.8)$$

To evaluate the upper bound of the right-hand side, each term is calculated using (I.2) and (I.5) as follows:

$$\int \varphi^2 dx = \int (f_n - f_m)^2 dx \leq 2 \int (f_n^2 + f_m^2) dx \leq 4C \quad (I.9)$$

$$\int \bar{\varphi}^2 dx = \int (\bar{f}_n - \bar{f}_m)^2 dx = \sum_i L_i \{\bar{f}_n(i) - \bar{f}_m(i)\}^2 \leq 2 \sum_i L_i \{\bar{f}_n^2(i) + \bar{f}_m^2(i)\} \quad (I.10)$$

For the last inequality, we used the relation $(A \pm B)^2 \leq 2(A^2 + B^2)$ which is equivalent to $0 \leq (A \pm B)^2$. From the definition of $\bar{f}_n(i)$, we have

$$\begin{aligned} \sum_i L_i \bar{f}_n^2(i) &= \sum_i L_i \left(L_i^{-1} \int_{L_i} f_n dx \right)^2 \leq \sum_i L_i (L_i^{-1})^2 L_i \int_{L_i} f_n^2 dx \\ &= \int_L f_n^2 dx \leq C \end{aligned} \quad (I.11)$$

where (I.1) and (I.5) are used. Combining (I.10) and (I.11), we have

$$\int \bar{\varphi}^2 dx \leq 4C \quad (I.12)$$

From (I.8), (I.9), and (I.12), we obtain

$$\int (\varphi + \bar{\varphi})^2 dx \leq 16C \quad (\text{I.13})$$

Similarly, using (I.4), we have

$$\begin{aligned} \int (\varphi - \bar{\varphi})^2 dx &\leq K \int (d\varphi/dx)^2 dx \\ &\leq 2K \int \{(df_n/dx)^2 + (df_m/dx)^2\} dx \leq 4KC \end{aligned} \quad (\text{I.14})$$

Substituting the right-hand sides of (I.13) and (I.14) in the square root of (I.7), we obtain

$$\int (\varphi^2 - \bar{\varphi}^2) dx \leq 8\sqrt{KC}$$

If $\max L_i \rightarrow 0$ as $N \rightarrow \infty$, then $K \rightarrow 0$ and (6) is proved.

In the second step, we prove the existence of a subsequence of the $\tilde{f}_n$'s such that

$$\lim_{m, n \rightarrow \infty} \int_L (\tilde{f}_n - \tilde{f}_m)^2 dx \rightarrow 0 \quad (\text{I.15})$$

A combination of (I.15) with (I.6) then gives Rellich's theorem.

We first note that

$$\begin{aligned} \int_L (\tilde{f}_n - \tilde{f}_m)^2 dx &= \sum_{i=1}^N L_i \{\tilde{f}_n(i) - \tilde{f}_m(i)\}^2 \\ &= \sum_{i=1}^N \{L_i^{1/2} \tilde{f}_n(i) - L_i^{1/2} \tilde{f}_m(i)\}^2 \end{aligned} \quad (\text{I.16})$$

and from (I.11)

$$L_i^{1/2} |\tilde{f}_n(i)| \leq \sqrt{C}$$

By Lemma 4, there is a subsequence of the $\tilde{f}_n$'s for which $L_1^{1/2} \tilde{f}_n(1)$ converges. Again by Lemma 4, using the members of this subsequence only, we can make another subsequence for which $L_2^{1/2} \tilde{f}_n(2)$ converges. Repeating this process, we obtain a subsequence of the $\tilde{f}_n$'s for which $L_i^{1/2} \tilde{f}_n(i)$ converges for every i . For this subsequence, the right-hand side of (I.16) approaches zero as $m, n \rightarrow \infty$, i.e., (I.15) is satisfied. The proof of Rellich's theorem is thus completed.

Proof for $\lim_{n \rightarrow \infty} k_n^2 = \infty$. Let E_{y_n} be the n th eigenfunction which is normalized then the corresponding eigenvalue k_n^2 can be expressed in the

form

$$k_n^2 = \int (dE_{yn}/dx)^2 dx$$

If k_n^2 remains finite, then the E_{yn} 's satisfy the condition for Rellich's theorem, and there must be E_{yn} and E_{ym} satisfying

$$\int (E_{yn} - E_{ym})^2 dx \leq \varepsilon$$

no matter how small ε is. However, the left-hand side is equal to

$$\int E_{ym}^2 dx - 2 \int E_{yn} E_{ym} dx + \int E_{yn}^2 dx$$

in which the first and third terms are both equal to 1, and the second term is equal to zero from the orthogonality condition. This means that $\int (E_{yn} - E_{ym})^2 dx$ cannot be less than 2 leading to a contradiction. We conclude from this that the assumption of k_n^2 remaining finite is wrong, or equivalently, that $\lim_{n \rightarrow \infty} k_n^2 = \infty$.

2. Two-Dimensional Case

LEMMA 1. Let f and g be single-valued real functions of x and y , then

$$\int f^2 dS \int g^2 dS \geq \left(\int fg dS \right)^2$$

where the integrals are over a certain area S .

LEMMA 2. Let f and g be single-valued real functions of x and y , then

$$\int (f \pm g)^2 dS \leq 2 \left(\int f^2 dS + \int g^2 dS \right)$$

where the integrals are over a certain area S .

The proofs of the above two lemmas are similar to those for the one-dimensional case, and hence will not be presented.

LEMMA 3. Let f be a single-valued real function of x and y having a square-integrable derivative ∇f over $S = \sum S_i$. Let $f(i)$ be the average of f over S_i ,

$$f(i) = S_i^{-1} \int_{S_i} f dS$$

and define the staircase function $\bar{f}$ whose value is equal to $f(i)$ in S_i for every i . We then have

$$\int_S (f - \bar{f})^2 dS \leq K \int_S (\nabla f)^2 dS$$

where K is a constant, and it is possible to choose the S_i 's so that $K \rightarrow 0$ when every $S_i \rightarrow 0$.

Proof. Suppose that S_i is an area enclosed by $x=0$, $x=a$, $y=0$, $y=g(x)$ such that the maximum value of $g(x)$ is smaller than h , and the minimum value is larger than t as shown in Fig. A.1. Let R be the shaded

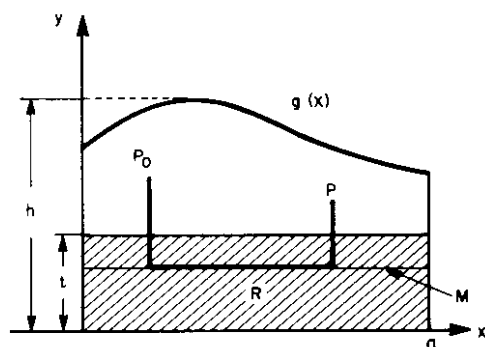


Fig. A.1. Domain S_i and the integral path $\overline{P_0P}$.

area below t , then there is a straight line M parallel to the x -axis and located in R such that

$$t \int_M (\nabla f)^2 dx = \iint_R (\nabla f)^2 dx dy \quad (\text{I.17})$$

Such a line M exists since the right-hand side can be written in the form

$$\int_0^t \int_0^a (\nabla f)^2 dx dy$$

where the integral with respect to x gives a continuous function of y which can be integrated from $y=0$ to t . There should be at least one point somewhere between $y=0$ and t at which the value of the continuous function of y becomes equal to the average value for this range of y . Let us draw a straight line M parallel to the x -axis through this point, then the left-hand

side of (I.17) becomes the average value times t , which is equal to the right-hand side.

Let $\mathbf{l}$ be a unit vector, then the value of

$$\nabla f \cdot \mathbf{l} = \partial f / \partial l$$

becomes maximum when the directions of $\mathbf{l}$ and ∇f coincide. Otherwise, the value becomes smaller than the magnitude of ∇f . It follows from this that

$$(\partial f / \partial l)^2 \leq \nabla f \cdot \nabla f = (\partial f / \partial x)^2 + (\partial f / \partial y)^2 \quad (\text{I.18})$$

Let P_0 be a point where f takes the value of $f(i)$ in S_i , then the value of f at P in S_i is given by

$$f = f(i) + \int_{P_0}^P (\partial f / \partial l) dl$$

from which we have

$$\begin{aligned} \{f - f(i)\}^2 &= \int_{P_0}^P (\partial f / \partial l) dl \int_{P_0}^P (\partial f / \partial l) dl \leq L \int_{P_0}^P (\partial f / \partial l)^2 dl \\ &\leq L \int_{P_0}^P \{(\partial f / \partial x)^2 + (\partial f / \partial y)^2\} dl = L \int_{P_0}^P (\nabla f)^2 dl \end{aligned}$$

where L is the path length from P_0 to P and the relations (I.1) and (I.18) are used. If we choose the path as shown in Fig. A.1, and integrate both sides over S_i the following inequality results.

$$\begin{aligned} \int_{S_i} \{f - f(i)\}^2 dS &\leq \int_{S_i} L \int_0^{P_0} (\nabla f)^2 dy dS + \int_{S_i} L \int_0^P (\nabla f)^2 dy dS \\ &\quad + \int_{S_i} L \int_M (\nabla f)^2 dx dS \end{aligned} \quad (\text{I.19})$$

Using the inequalities

$$\int_{S_i} \int_0^P (\nabla f)^2 dy dS \leq \int_0^h \int_0^a \int_0^{g(x)} (\nabla f)^2 dy dx dy \leq h \int_{S_i} (\nabla f)^2 dS$$

and

$$\int_{S_i} \int_M (\nabla f)^2 dx dS \leq \int_{S_i} t^{-1} \int_{S_i} (\nabla f)^2 dS dS \leq aht^{-1} \int_{S_i} (\nabla f)^2 dS$$

Eq. (I.19) becomes

$$\int_{S_i} \{f - f(i)\}^2 dS \leq (2h + a)(2h + aht^{-1}) \int_{S_i} (\nabla f)^2 dS \quad (\text{I.20})$$

where $L \leq 2h + a$ is used.

Let us divide the given area S into a large number of small areas, each of which satisfies the condition explained in connection with Fig. A.1; any rotation or translation of the coordinates will be allowed. Suppose that t of each area is kept equal to or larger than a constant α times h , i.e., $t \geq \alpha h$, then

$$\begin{aligned} \int_S (f - \bar{f})^2 dS &= \sum_i \int_{S_i} \{f - \bar{f}(i)\}^2 dS \\ &\leq \sum_i (2h + a)(2h + \alpha\alpha^{-1}) \int_{S_i} (\nabla f)^2 dS \leq K \int_S (\nabla f)^2 dS \end{aligned}$$

and K approaches zero as we decrease the maximum size of S_i . This completes the proof.

RELICH'S THEOREM. A sequence of functions $\bar{f}_n$ ($n=1, 2, \dots, \infty$), each satisfying

$$\int_S f_n^2 dS \leq C, \quad \int_S (\nabla f_n)^2 dS \leq C$$

contains a subsequence for which

$$\lim_{m, n \rightarrow \infty} \int_S (f_n - f_m)^2 dS = 0$$

The proof for the one-dimensional case will serve here almost word for word and hence we shall not repeat it.

Proof for $\lim_{n \rightarrow \infty} k_n^2 = \infty$. Let $\mathbf{E}_t$ be a normalized eigenfunction, then the corresponding eigenvalue k^2 can be expressed in the form

$$k^2 = \int (\nabla \times \mathbf{E}_t)^2 + (\nabla \cdot \mathbf{E}_t)^2 dS$$

which is equivalent to

$$\begin{aligned} k^2 &= \int \{ \mathbf{E}_t \cdot \nabla \times \nabla \times \mathbf{E}_t + \nabla \cdot (\mathbf{E}_t \times \nabla \times \mathbf{E}_t) + \nabla \cdot (\mathbf{E}_t \nabla \cdot \mathbf{E}_t) \\ &\quad - (\mathbf{E}_t \cdot \nabla) (\nabla \cdot \mathbf{E}_t) \} dS \end{aligned}$$

The integrals of the second and third terms in the brackets can be converted into line integrals, both of which disappear because of the boundary conditions for $\mathbf{E}_t$. The first term can be decomposed into two terms using the universal equality

$$\nabla \times \nabla \times \mathbf{E}_t = \nabla (\nabla \cdot \mathbf{E}_t) - \nabla^2 \mathbf{E}_t$$

the first of these, when multiplied by $\mathbf{E}_t$, is equal to the last term in the integrand. Thus, we are left with

$$\begin{aligned} k^2 &= - \int \mathbf{E}_t \cdot \nabla^2 \mathbf{E}_t \, dS \\ &= - \int \left(E_x \frac{\partial^2 E_x}{\partial x^2} + E_x \frac{\partial^2 E_x}{\partial y^2} + E_y \frac{\partial^2 E_y}{\partial x^2} + E_y \frac{\partial^2 E_y}{\partial y^2} \right) dS \end{aligned} \quad (\text{I.21})$$

Using

$$E_x \frac{\partial^2 E_x}{\partial x^2} = \frac{\partial}{\partial x} E_x \frac{\partial E_x}{\partial x} - \left(\frac{\partial E_x}{\partial x} \right)^2$$

and similar equations, (I.21) can be rewritten in the form

$$\begin{aligned} k^2 &= \int \{(\nabla E_x)^2 + (\nabla E_y)^2\} \, dS - \frac{1}{2} \int \left\{ \frac{\partial^2}{\partial x^2} (E_x^2 + E_y^2) + \frac{\partial^2}{\partial y^2} (E_x^2 + E_y^2) \right\} \, dS \\ &= \int \{(\nabla E_x)^2 + (\nabla E_y)^2\} \, dS - \frac{1}{2} \int \nabla \cdot \nabla (\mathbf{E}_t \cdot \mathbf{E}_t) \, dS \end{aligned} \quad (\text{I.22})$$

The second integral on the right-hand side can be converted into a line integral which disappears for the following reason. On the boundary $\nabla \cdot \mathbf{E}_t = 0$, and in terms of the $\mathbf{n}$ and $\mathbf{l}$ components this becomes

$$(\partial E_n / \partial n) + (\partial E_l / \partial l) = 0$$

Since E_l is equal to zero on the boundary, we have $(\partial E_l / \partial l) = 0$ and hence $(\partial E_n / \partial n) = 0$. The integrand of the line integral is given by

$$\mathbf{n} \cdot \nabla (\mathbf{E}_t \cdot \mathbf{E}_t) = \{\partial (E_n^2 + E_l^2) / \partial n\} = 2 \{E_n (\partial E_n / \partial n) + E_l (\partial E_l / \partial n)\}$$

which is equal to zero since $\partial E_n / \partial n$ and E_l both vanish on the boundary. Equation (I.22) now reduces to

$$k^2 = \int \{(\nabla E_x)^2 + (\nabla E_y)^2\} \, dS$$

Let E_{xn} and E_{yn} be the x and y components of the n th eigenfunction $\mathbf{E}_{in}$, then from the above expression of the eigenvalue, both the E_{xn} 's and the E_{yn} 's satisfy the conditions stated in Rellich's theorem if the k_n^2 's remain finite. Consequently, we can choose a subsequence of the $\mathbf{E}_{in}$'s for which the integral of $(E_{xn} - E_{xm})^2$ converges to zero as n and m grow indefinitely. Let us choose from this subsequence another subsequence for which the integral

of $(E_{yn} - E_{ym})^2$ converges to zero, then for such a subsequence, we have

$$\lim_{m, n \rightarrow \infty} \int (E_{in} - E_{im})^2 dS = \lim_{m, n \rightarrow \infty} \int \{(E_{xn} - E_{xm})^2 + (E_{yn} - E_{ym})^2\} dS = 0$$

However, this contradicts the orthogonality and normalization conditions for the E_{in} 's, which means that the assumption of the k_n^2 's remaining finite is wrong. This completes the proof of $\lim_{n \rightarrow \infty} k_n^2 = \infty$.

3. Three-Dimensional Case

Lemmas 1, 2, 3, and Rellich's theorem in the three-dimensional case are all identical to those in the two-dimensional case except that V replaces S everywhere. However, we have to use Fig. A.2 instead of Fig. A.1 to prove

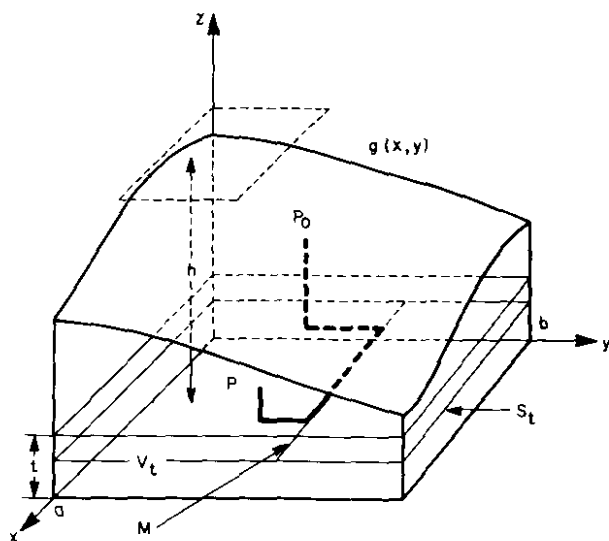


Fig. A.2. Domain V_t and the integral path $\overline{P_0P}$.

Lemma 3. Here, the maximum value of $g(x, y)$ is less than h and the minimum value is larger than $t > 0$. Also, V_t expresses the volume under t , and the plane S_t is such that

$$\int_{S_t} (\nabla f)^2 dS = t^{-1} \int_{V_t} (\nabla f)^2 dv$$

The line M on S_t is parallel to the x -axis and satisfies

$$\int_M (\nabla f)^2 dx = b^{-1} \int_{S_t} (\nabla f)^2 dS$$

If the integral path from P_0 to P is taken along the bold line shown in Fig. A.2, corresponding to (I.19) and (I.20) we have

$$\begin{aligned} \int_{v_i} \{f - f(i)\} dv &\leq L \left\{ \int_{v_i} \int_0^{P_0} (\nabla f)^2 dz dv \right. \\ &\quad + \int_{v_i} \int_0^P (\nabla f)^2 dz dv + \int_{v_i} \int_0^b (\nabla f)^2 dy dv \\ &\quad \left. + \int_{v_i} \int_0^b (\nabla f)^2 dy dv + \int_{v_i} \int_M (\nabla f)^2 dx dv \right\} \\ &\leq (2h + 2b + a) \{2h + 2bht^{-1} + abhb^{-1}t^{-1}\} \int_{v_i} (\nabla f)^2 dv \end{aligned}$$

If we divide the given volume V into a large number of smaller volumes similar to Fig. A.2 while keeping $t \geq \alpha h$, then

$$\int_V \{f - f(i)\} dv \leq K \int_V (\nabla f)^2 dv$$

and $K \rightarrow 0$ as the maximum size of the small volumes approaches zero. The proof of $\lim_{n \rightarrow \infty} k_n^2 = \infty$ is now obvious from the previous discussions for the one- and two-dimensional cases.

4. Resonant Cavities with Inhomogeneous Media

As we discussed in Section 4.5, the eigenfunctions for a resonant cavity with an inhomogeneous medium can be classified into three groups, I, II, and III. We shall first show that the number of independent functions in group I is equal to that of independent functions in group I for the same cavity with a homogeneous medium. We shall then show that the eigenvalue of the l th eigenfunction in group II cannot be smaller than that of the l th eigenfunction in group II for the same cavity with a homogeneous medium whose relative dielectric constant and permeability are given by ϵ_M and μ_M , respectively, where ϵ_M and μ_M are the maximum values of ϵ_r and μ_r in the inhomogeneous medium. Finally, we shall show that the eigenvalue of the l th eigenfunction in group III cannot be smaller than that of the l th eigen-

function in group III for the same cavity with a homogeneous medium whose relative dielectric constant and permeability are given by ε_m and μ_m respectively, where ε_m and μ_m are the minimum values of ε_r and μ_r . Since the eigenvalue k_n^2 for a cavity with a homogeneous medium grows indefinitely, k_n^2 for the cavity with an inhomogeneous medium must also grow indefinitely with n , which completes the proof.

Now suppose that $\mathbf{E}$ represents an eigenfunction in group I, then it satisfies $\nabla \times \mathbf{E} = 0$ in V and $\mathbf{n} \times \mathbf{E} = 0$ on S , where V indicates the volume in the cavity, and S the wall surface. By Helmholtz's theorem, $\mathbf{E}$ is expressible as $\nabla \varphi$. Substituting this into $\nabla \cdot \varepsilon_r \mathbf{E} = 0$ and $\mathbf{n} \times \mathbf{E} = 0$, we have

$$\nabla \cdot \varepsilon_r \nabla \varphi = 0 \quad (\text{in } V), \quad \mathbf{n} \times \nabla \varphi = 0 \quad (\text{on } S)$$

This boundary condition requires that φ is constant on each independent wall, and the problem is reduced to that of finding static potential functions in the closed region. Let N be the number of independent walls, then there are exactly $N-1$ independent solutions for φ since there are only $N-1$ independent ways of assigning the potential to $N-1$ walls with respect to the remaining one, whether or not the medium is homogeneous. This completes the discussion of group I.

Let us next consider $\mathbf{E}$ in group II. We shall discuss the problem in two steps. In the first step, we shall show a maximum-minimum property of eigenvalues. In the second step, we shall compare the magnitudes of the l th eigenvalues for the homogeneous and inhomogeneous cases.

For the functions in group II, $\nabla \cdot \varepsilon_r \mathbf{E} = 0$ and we have

$$\nabla \times \mu_r^{-1} \nabla \times \mathbf{E} - k^2 \varepsilon_r \mathbf{E} = 0 \quad (\text{in } V), \quad \mathbf{n} \times \mathbf{E} = 0 \quad (\text{on } S) \quad (\text{I.23})$$

If we restrict ourselves to the class of functions which satisfy $\mathbf{n} \times \mathbf{E} = 0$ on S and do not contain a term expressible as the gradient of a scalar function, the variational expression for the eigenvalue of (I.23) is given by

$$k^2(\mathbf{E}) = \frac{\int \mu_r^{-1} (\nabla \times \mathbf{E})^2 dv}{\int \varepsilon_r \mathbf{E}^2 dv}$$

The eigenfunctions in group III can all be expressed as the gradient of some scalar function, and hence they are excluded from the present $\mathbf{E}$'s. Let $\mathbf{F}_1, \mathbf{F}_2, \dots, \mathbf{F}_{l-1}$ be piecewise-continuous, but otherwise arbitrary functions,

and assume that $\mathbf{E}$ satisfies

$$\int \epsilon_r \mathbf{E} \cdot \mathbf{F}_i \, dv = 0 \quad (1 \leq i \leq l-1) \quad (\text{I.24})$$

$$\int \epsilon_r \mathbf{E}^2 \, dv = 1 \quad (\text{I.25})$$

Let us define $D(\mathbf{E})$ by

$$D(\mathbf{E}) = \int \mu_r^{-1} (\nabla \times \mathbf{E})^2 \, dv$$

then the value of $D(\mathbf{E})$ varies with $\mathbf{E}$. Let D be the minimum value of $D(\mathbf{E})$ when $\mathbf{E}$ varies over all possible functions satisfying (I.24) and (I.25) with the $\mathbf{F}_i$'s being fixed. If the $\mathbf{F}_i$'s change, D of course changes. Let us try to find the largest value of D when the $\mathbf{F}_i$'s are allowed to vary over all possible functions. We first consider a function in the form

$$\mathbf{E} = \sum_{i=1}^l c_i \mathbf{E}_i \quad (\text{I.26})$$

where $\mathbf{E}_1, \mathbf{E}_2, \dots, \mathbf{E}_l$ are the first l eigenfunctions in group II. From the conditions (I.24) and (I.25), all the c_i 's are determined. If $i \neq j$, the expression

$$\begin{aligned} \int \mu_r^{-1} (\nabla \times \mathbf{E}_i) \cdot (\nabla \times \mathbf{E}_j) \, dv &= \int \mathbf{E}_i \cdot \nabla \times \mu_r^{-1} \nabla \times \mathbf{E}_j \, dv \\ &+ \int \mathbf{E}_i \times \mu_r^{-1} \nabla \times \mathbf{E}_j \cdot \mathbf{n} \, dS \end{aligned}$$

is equal to zero, since the first term on the right-hand side is equal to $\int k_j^2 \epsilon_r \mathbf{E}_i \cdot \mathbf{E}_j \, dv$ which vanishes because of the orthogonality condition; the second term vanishes because $\mathbf{n} \times \mathbf{E}_i = 0$ on S . Using the above relation, the value of $D(\mathbf{E})$ for (I.26) can easily be calculated as

$$D(\mathbf{E}) = \sum_{i=1}^l c_i^2 k_i^2$$

However, since $k_i^2 \geq k_l^2$ ($1 \leq i \leq l-1$) and $\sum c_i^2 = 1$ from (I.25) and (I.26), we have

$$D(\mathbf{E}) \leq k_l^2$$

This inequality holds for at least one $\mathbf{E}$ given by (I.26) regardless of the $\mathbf{F}_i$'s used, which means $D \leq k_l^2$. Now suppose that $\mathbf{F}_1 = \mathbf{E}_1, \mathbf{F}_2 = \mathbf{E}_2, \dots, \mathbf{F}_{l-1} = \mathbf{E}_{l-1}$, then (I.24) gives the conditions under which the variational expression $k^2(\mathbf{E})$ is minimized to obtain $\mathbf{E}_l$. Because of (I.25), the minimum

value of $D(\mathbf{E})$ under these conditions is given by $k_l'^2$. Combining with $D \leq k_l'^2$, we conclude from the above discussion that the largest value of D is $k_l'^2$, and that the function which gives this value is the l th eigenfunction $\mathbf{E}_l$. This completes the discussion of the maximum-minimum property of the eigenvalues.

In the second step, we first change μ_r to μ_M . With the $\mathbf{F}_i$'s ($1 \leq i \leq l-1$) and $\mathbf{E}$ being fixed,

$$D'(\mathbf{E}) = \int \mu_M^{-1} (\nabla \times \mathbf{E})^2 dS$$

cannot be larger than $D(\mathbf{E})$. The minimum value D' of $D'(\mathbf{E})$ cannot therefore be larger than D . It follows from this that the l th eigenvalue $k_l'^2$ in group II for the same cavity with ϵ_r and μ_M is smaller than $k_l'^2$, thus

$$k_l'^2 \leq k_l'^2$$

because the largest value of D' gives $k_l'^2$. Next, we change ϵ_r to ϵ_M . The conditions necessary to obtain the l th eigenvalue $k_l''^2$ for the cavity with ϵ_M and μ_M become

$$\int \epsilon_M \mathbf{E} \cdot \mathbf{F}_i' dv = 0 \quad (1 \leq i \leq l-1) \quad (\text{I.27})$$

$$\int \epsilon_M \mathbf{E}^2 dv = 1 \quad (\text{I.28})$$

If we set $\mathbf{F}_i' = \mathbf{F}_i(\epsilon_r/\epsilon_M)$, (I.27) reduces to (I.24), in order to satisfy (I.28), however, $\mathbf{E}$ must be multiplied by a constant c smaller than unity. The corresponding $D(\mathbf{E})$, which we shall indicate by $D''(\mathbf{E})$, must be equal to $c^2 D'(\mathbf{E})$, and hence

$$D''(\mathbf{E}) \leq D'(\mathbf{E})$$

Let D'' be the minimum value of $D''(\mathbf{E})$ when $\mathbf{E}$ varies over all possible functions, then $D'' \leq D'$. If we change $\mathbf{F}_i$ over all possible functions, the corresponding $\mathbf{F}_i'$ also varies over all possible functions. However, the largest value of D'' cannot exceed the largest value of D' because of the relation $D'' \leq D'$, which means $k_l''^2 \leq k_l'^2 \leq k_l'^2$. We conclude from this that the eigenvalue of the l th eigenfunction in group II, for a cavity with ϵ_r and μ_r , cannot be smaller than that of the l th eigenfunction in group II for the same cavity with ϵ_M and μ_M . Since the latter grows indefinitely with l , so does the former.

We finally consider $\mathbf{E}$ in group III. Since $\nabla \times \mathbf{E} = 0$ in V , the eigenvalue

problem for $\mathbf{E}$ is given by

$$\nabla \nabla \cdot \varepsilon_r \mathbf{E} + k^2 \mathbf{E} = 0 \quad (\text{in } V), \quad \mathbf{n} \times \mathbf{E} = 0, \quad \nabla \cdot \varepsilon_r \mathbf{E} = 0 \quad (\text{on } S)$$

Writing $\nabla \cdot \varepsilon_r \mathbf{E}$ as $k\varphi$, we obtain

$$\nabla \cdot \varepsilon_r \nabla \varphi + k^2 \varphi = 0 \quad (\text{in } V), \quad \varphi = 0 \quad (\text{on } S) \quad (\text{I.29})$$

For each φ satisfying (I.29), a corresponding $\mathbf{E}$ can be obtained from

$$k\mathbf{E} = -\nabla \varphi$$

Note that $\mathbf{n} \times \mathbf{E} = 0$ on S is automatically satisfied. If we restrict ourselves to the functions which satisfy $\varphi = 0$ on S , the variational expression for the eigenvalues of (I.29) is given by

$$k^2(\varphi) = \frac{\int \varepsilon_r (\nabla \varphi)^2 dv}{\int \varphi^2 dv}$$

Let $\psi_1, \psi_2, \dots, \psi_{l-1}$ be piecewise-continuous, but otherwise arbitrary, functions and assume that φ satisfies

$$\int \varphi \psi_i dv = 0 \quad (1 \leq i \leq l-1) \quad (\text{I.30})$$

$$\int \varphi^2 dv = 1 \quad (\text{I.31})$$

Let us define $D(\varphi)$ by

$$D(\varphi) = \int \varepsilon_r (\nabla \varphi)^2 dv$$

and let D be the minimum value of $D(\varphi)$ when φ varies over all possible functions, then an argument similar to the one employed for group II shows that the eigenvalue of the l th eigenfunction in group III for a cavity with ε_r cannot be smaller than that for the same cavity with ε_m . This completes the proof for $\mathbf{E}$.

For $\mathbf{H}$, the number of independent eigenfunctions in group I is equal to the number of independent loops formed by the wall surfaces. Apart from this, all the arguments are similar to those for $\mathbf{E}$.

APPENDIX II

GENERALIZED FUNCTIONS AND FOURIER ANALYSIS

In this appendix, we shall introduce the concept of generalized functions in order to simplify our discussion of oscillator noise in Section 9.3. This concept was developed from Dirac's "delta function" $\delta(x)$ which satisfies

$$\int_{-\infty}^{\infty} \delta(x) K(x) dx = K(0)$$

for any "reasonable" function $K(x)$. No ordinary function can satisfy the above equation. However, we can consider the limit of a sequence of functions such as $(n/\pi)^{1/2} \exp(-nx^2)$ ($n=1, 2, 3, \dots$). The value of these functions tends to be zero with increasing n at almost every point except $x=0$ while the area under the curves described by the functions remains equal to unity. The following discussion will give this intuitive approach a mathematical foundation, and not only justify the use of $\delta(x)$ but also greatly reduce the complications in Fourier analysis such as the determination of conditions for convergence, term by term differentiability, and uniqueness.

We start with terminology. A function is said to be good if it is everywhere differentiable any number of times; also, its values and those of all its derivatives reach the order of $|x|^{-N}$ at most as $|x|$ increases, i.e., $O(|x|^{-N})$, for all N . A sequence $f_n(x)$ of good functions defines a generalized function $f(x)$, if

$$\lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f_n(x) K(x) dx$$

exists for any good function $K(x)$. The integral of the product of a generalized

function $f(x)$ and a good function $K(x)$ is defined by

$$\int_{-\infty}^{\infty} f(x) K(x) dx = \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f_n(x) K(x) dx \quad (\text{II.1})$$

If two sequences of good functions, $f_n(x)$ and $g_n(x)$, satisfy

$$\lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f_n(x) K(x) dx = \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} g_n(x) K(x) dx \quad (\text{II.2})$$

for any good function $K(x)$, then the sequences are said to be equivalent, and the corresponding generalized functions $f(x)$ and $g(x)$ are said to be equal. The equality is indicated by

$$f(x) = g(x) \quad (\text{II.3})$$

Thus, each generalized function is really the class of sequences all equivalent to each other. The above definition of equality is, of course, consistent with the definition of integral (II.1) since the limit on the right-hand side is the same for all equivalent sequences (see p. 425 for the physical meaning).

If two generalized functions $f(x)$ and $g(x)$ are defined by two sequences $f_n(x)$ and $g_n(x)$, then the sum $f(x) + g(x)$ is defined by the sequence $f_n(x) + g_n(x)$. Also, a generalized function $f(ax + b)$ is defined by the sequence $f_n(ax + b)$, where a and b are constants. The derivative $f'(x)$ of a generalized function is defined by the sequence $f'_n(x)$. Since $K'(x)$ is a good function for any good function $K(x)$, the limit

$$\lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f'_n(x) K(x) dx = - \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f_n(x) K'(x) dx \quad (\text{II.4})$$

exists and is the same for all the sequences equivalent to $f_n(x)$. From (II.4), we have

$$\int_{-\infty}^{\infty} f'(x) K(x) dx = - \int_{-\infty}^{\infty} f(x) K'(x) dx \quad (\text{II.5})$$

which is consistent with the definition of the equality between generalized functions. If two generalized functions $g(x)$ and $f(x)$ are equal, $f'(x) = g'(x)$.

Let us next consider Fourier transformations. The Fourier transform of a good function $f_n(x)$,

$$F_n(y) = \int_{-\infty}^{\infty} f_n(x) e^{-j2\pi xy} dx \quad (\text{II.6})$$

is also a good function. This can be shown by differentiating (II.6) m times

and integrating by parts N times as follows:

$$\begin{aligned} \left| \frac{d^m}{dy^m} F_n(y) \right| &= \left| \frac{1}{(j2\pi y)^N} \int_{-\infty}^{\infty} \frac{d^N}{dx^N} \{(-j2\pi x)^m f_n(x)\} e^{-j2\pi xy} dx \right| \\ &\leq \frac{(2\pi)^{m-N}}{|y|^N} \int_{-\infty}^{\infty} \left| \frac{d^N}{dx^N} \{x^m f_n(x)\} \right| dx = O(|y|^{-N}) \end{aligned}$$

The inverse transform of (II.6) is given by

$$f_n(x) = \int_{-\infty}^{\infty} F_n(y) e^{j2\pi xy} dy \quad (\text{II.7})$$

To prove this, let us investigate the magnitude of

$$f_n(x) - \int_{-\infty}^{\infty} F_n(y) \exp(-\delta^2 y^2) e^{j2\pi xy} dy$$

Substituting (II.6) into this expression and noting that both $f_n(x)$ and $F_n(y)$ are good functions, we have

$$\begin{aligned} &\left| f_n(x) - \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f_n(t) \exp(-\delta^2 y^2) e^{j2\pi(x-t)y} dt dy \right| \\ &\leq \left| f_n(x) - \frac{\sqrt{\pi}}{\delta} \int_{-\infty}^{\infty} f_n(t) \exp\left\{-\left(\frac{\pi}{\delta}\right)^2 (t-x)^2\right\} dt \right| \\ &= \left| \frac{1}{\sqrt{\pi}} \int_{-\infty}^{\infty} f_n(x) \exp(-z^2) dz - \frac{1}{\sqrt{\pi}} \int_{-\infty}^{\infty} f_n\left(x + \frac{\delta}{\pi} z\right) \exp(-z^2) dz \right| \\ &\leq \frac{1}{\sqrt{\pi}} \int_{-\infty}^{\infty} \max |f'_n(x)| \left| \frac{\delta}{\pi} z \right| \exp(-z^2) dz \quad (\text{II.8}) \end{aligned}$$

which leads to (II.7) in the limit of $\delta \rightarrow 0$. In the above calculation, we have used the relation

$$\int_{-\infty}^{\infty} \exp(-z^2) dz = \sqrt{\pi} \quad (\text{II.9})$$

which can be derived using a transformation to a polar coordinate as follows:

$$\begin{aligned} \left(\int_{-\infty}^{\infty} \exp(-z^2) dz \right)^2 &= \int_{-\infty}^{\infty} \exp(-x^2) dx \int_{-\infty}^{\infty} \exp(-y^2) dy \\ &= \int_0^{\infty} \exp(-r^2) 2\pi r dr = \pi \end{aligned}$$

Notice that the limits of the integral on the right-hand side of (II.9) can be

replaced by $-\infty + jA$ and $\infty + jA$ without changing the value of the integral where A is a real constant.

Let us now define the Fourier transform $F(y)$ of a generalized function $f(x)$ by the sequence $F_n(y)$, where $F_n(y)$ is the Fourier transform of $f_n(x)$, then

$$F(y) = \int_{-\infty}^{\infty} f(x) e^{-j2\pi xy} dx \quad (\text{II.10})$$

while the inverse transform of $F(y)$ becomes $f(x)$ from (II.7):

$$f(x) = \int_{-\infty}^{\infty} F(y) e^{j2\pi xy} dy \quad (\text{II.11})$$

Notice that the Fourier transforms of $f(ax + b)$ and $f'(x)$ are given by $|a|^{-1} e^{j2\pi by/a} F(y/a)$ and $j2\pi y F(y)$, respectively. Furthermore, since

$$\begin{aligned} \int_{-\infty}^{\infty} f_n(x) K(x) dx &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} F_n(y) K(x) e^{j2\pi xy} dx dy \\ &= \int_{-\infty}^{\infty} F_n(y) k(y) dy \end{aligned} \quad (\text{II.12})$$

we have

$$\int_{-\infty}^{\infty} f(x) K(x) dx = \int_{-\infty}^{\infty} F(y) k(y) dy \quad (\text{II.13})$$

where $k(y)$ is the inverse transform of $K(x)$.

We are now in a position to define a generalized function $\delta(x)$. Let us consider a sequence $(n/\pi)^{1/2} \exp(-nx^2)$. First, we note that each member of the sequence is a good function. Next, since

$$\begin{aligned} &\left| \int_{-\infty}^{\infty} (n/\pi)^{1/2} \exp(-nx^2) K(x) dx - K(0) \right| \\ &= \left| \int_{-\infty}^{\infty} (n/\pi)^{1/2} \exp(-nx^2) \{K(x) - K(0)\} dx \right| \\ &\leq \max |K'(x)| \int_{-\infty}^{\infty} (n/\pi)^{1/2} \exp(-nx^2) |x| dx \\ &= (\pi n)^{-1/2} \max |K'(x)| \rightarrow 0 \end{aligned}$$

as $n \rightarrow \infty$, we have

$$\lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} (n/\pi)^{1/2} \exp(-n\pi^2 x^2) K(x) dx = K(0)$$

for any good function $K(x)$. Thus, sequences equivalent to $(n/\pi)^{1/2} \exp(-nx^2)$

define a generalized function. If we indicate this generalized function by $\delta(x)$, we have

$$\int_{-\infty}^{\infty} \delta(x) K(x) dx = K(0) \quad (\text{II.14})$$

Since the Fourier transform of $(n/\pi)^{1/2} \exp(-nx^2)$ is $\exp(-\pi^2 y^2/n)$, the Fourier transform of $\delta(x)$ is one.

An ordinary function $f_0(x)$ can be considered as a generalized function if $(1+x^2)^{-N} f_0(x)$ is absolutely integrable from $-\infty$ to ∞ . To show this, let us consider a sequence defined by

$$f_n(x) = \int_{-\infty}^{\infty} f_0(t) \sigma(nt - nx) n \exp(-t^2/n^2) dt \quad (\text{II.15})$$

where $\sigma(z)$ is any good function which is zero for $|z| \geq 1$ and positive for $|z| < 1$ while $\int_{-1}^1 \sigma(z) dz = 1$. Such a function is obtained, for example, by setting $\alpha = -1$ and $\beta = 1$ in $\psi_{\alpha\beta}(z)$ given by

$$\psi_{\alpha\beta}(z) = A^{-1} \exp \left\{ \frac{-1}{(z-\alpha)^2} + \frac{-1}{(z-\beta)^2} \right\} \quad \text{for } \alpha < z < \beta$$

$$= 0 \quad \text{for } z \leq \alpha \quad \text{or} \quad z \geq \beta \quad (\text{II.16})$$

where

$$A = \int_{\alpha}^{\beta} \exp \left\{ \frac{-1}{(z-\alpha)^2} + \frac{-1}{(z-\beta)^2} \right\} dz$$

Note that all the derivatives of $\psi_{\alpha\beta}(z)$ exist and are equal to zero if $z \leq \alpha$ or $z \geq \beta$. In order to show that $f_n(x)$ is a good function, we differentiate $f_n(x)$ m times. Since $\sigma^{(m)}(nt - nx)$ is nonzero only when $|x| - 1 < |t| < |x| + 1$, we have

$$\left| \frac{d^m}{dx^m} f_n(x) \right| = \int_{-\infty}^{\infty} f_0(t) (-n)^m \sigma^{(m)}(nt - nx) n \exp(-t^2/n^2) dt$$

$$\leq n^{m+1} \max |\sigma^{(m)}(z)| \exp \{ -(|x| - 1)^2/n^2 \}$$

$$\times \{ 1 + (|x| + 1)^2 \}^N \int_{-\infty}^{\infty} (1 + t^2)^{-N} |f_0(t)| dt = O(|x|^{-M})$$

as $|x| \rightarrow \infty$ for all M . This shows that $f_n(x)$ is, indeed, a good function. Next, let us consider

$$\left| \int_{-\infty}^{\infty} f_n(x) K(x) dx - \int_{-\infty}^{\infty} f_0(x) K(x) dx \right|$$

where $K(x)$ is an arbitrary good function. Substituting (II.15) into this expression, we have

$$\begin{aligned} & \left| \int_{-1}^1 \sigma(y) dy \int_{-\infty}^{\infty} f_0(t) \exp\left(-\frac{t^2}{n^2}\right) K\left(t - \frac{y}{n}\right) dt - \int_{-\infty}^{\infty} f_0(x) K(x) dx \right| \\ & \leq \max_{|y| < 1} \left| \int_{-\infty}^{\infty} f_0(t) \exp\left(-\frac{t^2}{n^2}\right) \left\{ K\left(t - \frac{y}{n}\right) - K(t) \right\} dt \right. \\ & \quad \left. - \int_{-\infty}^{\infty} f_0(t) K(t) \left\{ 1 - \exp\left(-\frac{t^2}{n^2}\right) \right\} dt \right| \\ & \leq \int_{-\infty}^{\infty} |f_0(t)| \left\{ \frac{1}{n} \max_{|x-t| < 1} |K'(x)| \right\} dt + \int_{-\infty}^{\infty} |f_0(t) K(t)| \frac{t^2}{n^2} dt \\ & \leq \frac{1}{n} \int_{-\infty}^{\infty} |f_0(t)| \frac{A}{(1+t^2)^N} dt + \frac{1}{n^2} \int_{-\infty}^{\infty} |f_0(t)| \frac{B}{(1+t^2)^N} dt \\ & = O\left(\frac{1}{n}\right) \end{aligned}$$

as $n \rightarrow \infty$, where A and B are constants. The above calculation shows that

$$\lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f_n(x) K(x) dx = \int_{-\infty}^{\infty} f_0(x) K(x) dx$$

for any good function $K(x)$. Consequently, the sequence $f_n(x)$ defines a generalized function $f(x)$ and $f(x)$ can be considered equal to $f_0(x)$ from the definition of equality between two generalized functions. Furthermore, the derivative and Fourier transform of $f(x)$ can be considered equal to those of $f_0(x)$ under certain conditions. Let us assume that $(1+|x|^2)^{-N} f_0'(x)$ is also absolutely integrable from $-\infty$ to $+\infty$ for some N , then we have

$$\int_{-\infty}^{\infty} f_0'(x) K(x) dx = - \int_{-\infty}^{\infty} f_0(x) K'(x) dx$$

A comparison of this relation with (II.5) shows that

$$f'(x) = f_0'(x)$$

in the sense of equality between two generalized functions. To show that the Fourier transform $F(y)$ of $f(x)$ is equal to the Fourier transform $F_0(y)$ of $f_0(x)$ when $F_0(y)$ exists, we have only to prove that

$$\int_{-\infty}^{\infty} F(y) k(y) dy = \int_{-\infty}^{\infty} F_0(y) k(y) dy$$

holds for any good function $k(y)$. This is obvious from the following calculation:

$$\begin{aligned}\int_{-\infty}^{\infty} F(y) k(y) dy &= \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} F_n(y) k(y) dy \\ &= \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f_n(x) K(x) dx = \int_{-\infty}^{\infty} f_0(x) K(x) dx \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} F_0(y) e^{j2\pi xy} dy K(x) dx = \int_{-\infty}^{\infty} F_0(y) k(y) dy\end{aligned}$$

where $K(x)$ is the Fourier transform of $k(y)$ and use is made of (II.12) for the second equality.

The above discussion shows for most calculations of interest that it is not necessary to distinguish between an ordinary function and the corresponding generalized function defined by (II.15). We shall henceforth omit the subscript 0 for distinguishing ordinary functions from generalized ones.

Let $f(x)$ and $g(x)$ be generalized functions defined by sequences of good functions $f_n(x)$ and $g_n(x)$, respectively, such that

$$h_n(x) = \int_{-\infty}^{\infty} f(t) g_n(x-t) dt$$

exists for all n , and that the sequence $h_n(x)$ defines a generalized function $h(x)$. This $h(x)$ is called the convolution of $f(x)$ and $g(x)$ and indicated by

$$h(x) = \int_{-\infty}^{\infty} f(t) g(x-t) dt \quad (\text{II.17})$$

The integrand is the product of two generalized functions. It is worth noting that if we arbitrarily select two generalized functions, their product may not necessarily be defined. This is because there is no guarantee that the limit

$$\lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f_n(x) g_n(x) K(x) dx$$

exists for every good function $K(x)$ even though $f_n(x) g_n(x)$ is a good function.

The Fourier transform $H(y)$ of $h(x)$ is given by the product of the Fourier transforms of $f(x)$ and $g(x)$. This can be seen as follows:

$$H(y) = \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(t) g_n(x-t) dt e^{-j2\pi xy} dx$$

$$\begin{aligned}
&= \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(t) g_n(\xi) e^{-j2\pi(t+\xi)y} dt d\xi \\
&= \lim_{n \rightarrow \infty} \int_{-\infty}^{\infty} f(t) e^{-j2\pi ty} dt \int_{-\infty}^{\infty} g_n(\xi) e^{-j2\pi \xi y} d\xi = F(y) G(y) \quad (\text{II.18})
\end{aligned}$$

If $f(x) = \delta(x)$ and the sequence $g_n(x)$ defines a generalized function $g(x)$, then the sequence

$$h_n(x) = \int_{-\infty}^{\infty} \delta(t) g_n(x-t) dt = g_n(x)$$

defines the same generalized function $g(x)$, and we have

$$g(x) = \int_{-\infty}^{\infty} \delta(t) g(x-t) dt \quad (\text{II.19})$$

Since the Fourier transform of $\delta(x)$ is 1, Eq. (II.18) gives a trivial result $G(y) = G(y)$ in this case.

If we set $g(x) = \delta(x)$ in (II.19) and use $\delta(x) = \delta(-x)$, we have

$$\delta(\tau) = \int_{-\infty}^{\infty} \delta(t) \delta(t+\tau) dt \quad (\text{II.20})$$

Similarly, if we set $g(x)$ equal to 1 from $-\theta$ to θ and zero otherwise, we have

$$g(0) = 1 = \int_{-\theta}^{\theta} \delta(t) dt \quad (\text{II.21})$$

where θ is a nonzero positive quantity.

If $f_t(x)$ is a generalized function for each value of the parameter t and $f(x)$ is another generalized function such that for any good function $K(x)$

$$\lim_{t \rightarrow c} \int_{-\infty}^{\infty} f_t(x) K(x) dx = \int_{-\infty}^{\infty} f(x) K(x) dx \quad (\text{II.22})$$

then we say

$$\lim_{t \rightarrow c} f_t(x) = f(x) \quad (\text{II.23})$$

Here c may be finite or infinite. Let $F_t(y)$ and $F(y)$ be the Fourier transforms of $f_t(x)$ and $f(x)$, then

$$\begin{aligned}
\int_{-\infty}^{\infty} F(y) k(y) dy &= \int_{-\infty}^{\infty} f(x) K(x) dx \\
&= \lim_{t \rightarrow c} \int_{-\infty}^{\infty} f_t(x) K(x) dx = \lim_{t \rightarrow c} \int_{-\infty}^{\infty} F_t(y) k(y) dy
\end{aligned}$$

or equivalently

$$\lim_{t \rightarrow c} F_t(y) = F(y) \quad (\text{II.24})$$

In other words, the Fourier transform of the limit of generalized functions is equal to the limit of the Fourier transforms of the generalized functions.

Let $x(t)$ be a function of time such that a generalized function $x_T(t)$ can be defined by

$$\begin{aligned} x_T(t) &= x(t), & -T < t < T \\ &= 0, & t \geq |T| \end{aligned} \quad (\text{II.25})$$

then the autocorrelation function of $x(t)$ is given by

$$R_x(\tau) = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-\infty}^{\infty} x_T(t) x_T(t + \tau) dt \quad (\text{II.26})$$

Note that $R_x(\tau)$ is an even function of τ since a change in the sign of τ does not change the value of the integral on the right-hand side. The Fourier transform of $x_T(x)$ is given by

$$X_T(f) = \int_{-\infty}^{\infty} x_T(t) e^{-j2\pi f t} dt \quad (\text{II.27})$$

from which we have

$$X_T(f) X_T^*(f) = \int_{-\infty}^{\infty} x_T(\theta) e^{-j2\pi f \theta} d\theta \int_{-\infty}^{\infty} x_T(t) e^{j2\pi f t} dt$$

Multiplying both sides by $e^{j2\pi f \tau}$ and integrating from $-\infty$ to ∞ with respect to f , we obtain

$$\begin{aligned} \int_{-\infty}^{\infty} |X_T(f)|^2 e^{j2\pi f \tau} df &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_T(t) x_T(\theta) e^{j2\pi f(\tau + t - \theta)} dt d\theta df \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_T(t) x_T(\theta) \delta(\theta - t - \tau) d\theta dt \\ &= \int_{-\infty}^{\infty} x_T(t) x_T(t + \tau) dt \end{aligned}$$

Dividing by $2T$ and taking the limit of $T \rightarrow \infty$, a comparison with (II.26) shows that

$$R_x(\tau) = \int_{-\infty}^{\infty} |x(f)|^2 e^{j2\pi f \tau} df \quad (\text{II.28})$$

where

$$|x(f)|^2 = \lim_{T \rightarrow \infty} \frac{1}{2T} |X_T(f)|^2 \quad (\text{II.29})$$

The function $|x(f)|^2$ is called the power spectral density of $x(t)$. The inverse transform of (II.28) gives

$$|x(f)|^2 = \int_{-\infty}^{\infty} R_x(\tau) e^{-j2\pi f\tau} d\tau \quad (\text{II.30})$$

Since $R_x(\tau)$ is an even function of τ , (II.30) can be rewritten as

$$|x(f)|^2 = 2 \int_0^{\infty} R_x(\tau) \cos 2\pi f\tau d\tau$$

Similarly, since $|x(f)|^2$ is an even function of f from the above equation, (II.28) can be written in the form

$$R_x(\tau) = 2 \int_0^{\infty} |x(f)|^2 \cos 2\pi f\tau df$$

If $x(f)$ represents a stationary random process which is ergodic (time averages and ensemble averages are equal), then $R_x(\tau)$ can also be written in the form

$$R_x(\tau) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_1 x_2 p(x_1) p(x_2 | x_1; \tau) dx_1 dx_2 \quad (\text{II.31})$$

where x_1 and x_2 indicate the values of $x(t)$ at times τ apart, $p(x_1)$ is the probability density function that $x(t)$ is in the vicinity of x_1 , and $p(x_2 | x_1; \tau)$ is the conditional probability density function that $x(t)$ is in the vicinity of x_2 at a time τ after it was in the vicinity of x_1 .

In much the same way, the cross-correlation function between two functions $x(t)$ and $y(t)$ is defined by

$$R_{xy}(\tau) = \lim_{T \rightarrow \infty} \frac{1}{2T} \int_{-T}^T x_T(t) y_T(t + \tau) dt \quad (\text{II.32})$$

and its Fourier transform is called the cross-power spectral density. When $x(t)$ and $y(t)$ represent stationary random processes which are jointly ergodic, $R_{xy}(\tau)$ can be written in the form

$$R_{xy}(\tau) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_1 y_1 p(x_1) p(y_1 | x_1; \tau) dx_1 dy_1 \quad (\text{II.33})$$

where $p(y_1 | x_1; \tau)$ is the conditional probability density function that $y(t)$ is in the vicinity y_1 at a time τ after $x(t)$ was in the vicinity of x_1 .

One of the most important concepts introduced in this appendix is that of equality between generalized functions. Whenever we have a function

representing an impulse such as a hammer blow or a surge of electric current, we are not interested in the detailed shape of the function but rather the consequence of such a blow applied to some physical system, i.e., the consequence which can generally be expressed as the integral of the impulse function with a proper weighting function. As a result, if two impulse functions give identical results for all possible weighting functions, they can be considered equal. This is exactly the concept used to define the equality between generalized functions.

Our mathematical equations are always idealizations for describing physical processes. The equality is examined by measurements which will never be able to test the values of functions at every point of space or at every instant of time. Instead, the examination will be made between the weighted averages over a small area in the vicinity of the point or over a small time or frequency interval. Thus, the equality we defined for generalized functions is more natural and more appropriate in the discussion of physical processes than the usual equality which demands the same values at each point of space or time. Looking again at Section 9.3, we can see how natural the noise discussion can be made with the help of generalized functions. The application of generalized functions is not restricted to the discussion of noise; there are numerous fields in which generalized functions prove to be useful although we did not discuss them in this book.

REFERENCES

The reader may wish to supplement his reading by referring to original papers which introduce some of the ideas discussed in this book; these will be given here.

The Smith chart was originally introduced by

SMITH, P. H., Transmission Line Calculator. *Electronics* 12, 29-31 (1939).

The equivalent circuit of two-port networks discussed in Section 1.3 is due to

WEISSFLOCH, A., Kreisgeometrische Vierpoltheorie und ihre Bedeutung für Messtechnik und Schaltungstheorie des Dezimeter- und Zentimeterwellengebietes. *Hochfrequenz. u. Elek.* 61, 100-123 (1943).

Power waves were first used by

PENFIELD, JR., P., Noise in Negative Resistance Amplifiers. *IRE Trans.* CT-7, 166-170 (1960).

A detailed discussion has been presented by

KUROKAWA, K., Power Waves and The Scattering Matrix. *IEEE Trans.* MTT-13, 194-202 (1965).

For proof of Helmholtz's theorem, the author consulted

MORSE, P. M., and FESHBACH, H., "Methods of Theoretical Physics." McGraw-Hill, New York, 1953.

For supplementing Chapter 2, the following two books are highly recommended:

STRATTON, J. A., "Electromagnetic Theory." McGraw-Hill, New York, 1941.

PANOFSKY, W. K. H., and PHILLIPS, M., "Classical Electricity and Magnetism." Addison-Wesley, Reading, Massachusetts, 1955.

A theory of waveguides was first discussed by

STRUTT (Lord Rayleigh), J. W., On the Passage of Electric Waves through Tubes or the Vibrations of Dielectric Cylinders. *Phil. Mag.* 43, 125-132 (1897).

Most of the fundamental concepts of waveguides, such as modes and cutoff frequencies, were established in this paper (Strutt); however, our discussion is largely based on

SLATER, J. C., "Microwave Electronics." Van Nostrand, Princeton, New Jersey, 1950.

This is the best book that the present author has read in the field of microwave circuits, and his understanding of microwave circuits was greatly influenced by it.

A good reference book for eigenvalue problems is

COURANT, R., and HILBERT, D., "Methods of Mathematical Physics," Vol. I. Wiley (Interscience), New York, 1953.

The discussion of Bessel functions has been restricted to a minimum in the preceding chapters; those who wish to learn more about these functions may begin with

McLACHLAN, N. W., "Bessel Functions for Engineers." Oxford Univ. Press (Clarendon), London and New York, 1934.

The first paper to discuss the effect of wall losses in waveguides using a perturbation method is

PAPADOPOULUS, V. M., Propagation of Electromagnetic Waves in Cylindrical Waveguides with Imperfectly Conducting Walls. *Quart. J. Mech. Appl. Math.* **7**, 325-334 (1954).

The set of functions which Slater used in his book for the series expansion of the magnetic field in cavities was found to be incomplete by

TEICHMANN, T., and WIGNER, E. P., Electromagnetic Field Expansions in Loss-free Cavities Excited through Holes. *J. Appl. Phys.* **24**, 262-267 (1953)

The discussion of resonant cavities in Section 4.2 is based on

KUROKAWA, K., The Expansions of Electromagnetic Fields in Cavities. *IRE Trans. MTT-6*, 178-187 (1958).

A detailed discussion of waveguide junctions is presented by

MONTGOMERY, C. G., DICKE, R. H., and PURCELL, E. M., "Principles of Microwave Circuits." McGraw-Hill, New York, 1948.

This is probably the first book devoted to the theory of microwave circuits.

Nonreciprocal circuits using ferrites were introduced in the microwave field by

HOGAN, C. L., The Ferromagnetic Faraday Effect at Microwave Frequencies and Its Applications--The Microwave Gyrator. *Bell System Tech. J.* **31**, 1-31 (1952).

The theory of Y -junction circulators in Section 5.8 is originally based on

AULD, B. A., The Synthesis of Symmetrical Waveguide Circulators. *IRE Trans. MTT-7*, 238-246 (1959).

The discussion of distributed couplings in Section 6.1 is based on

HAUS, H. A., Coupling of Modes of Propagation. MIT, Cambridge, Massachusetts.

Internal memorandum, copies of which are no longer available.

A concise discussion of periodic structures can be found in

WATKINS, D. A., "Topic in Electromagnetic Theory." Wiley, New York, 1958.

The canonical form of noisy n -port networks was published by

HAUS, H. A., and ADLER, R. B., Canonical Form of Linear Noisy Networks. *IRE Trans. CT-5*, 161-167 (1958).

The same theory is also presented by

HAUS, H. A., and ADLER, R. B., "Circuit Theory of Linear Noisy Networks." Wiley, New York, 1959.

These authors used the exchangeable noise measure for their discussion of amplifier noise performance. The noise measure, which includes the noise contribution originating in and reflected back to the load, was first introduced by

KUROKAWA, K., Actual Noise Measure of Linear Amplifiers. *Proc. IRE* **49**, 1391-1397 (1961).

The operating noise temperature, which includes the load noise contribution, is defined in IRE Standards on Electron Tubes issued in *Proc. IEEE* (March 1963).

Statistical noise theory has not been discussed in this book; those who wish to acquire background information are referred to

DAVENPORT, Jr., W. B., and ROOT, W. L., "Random Signals and Noise." McGraw-Hill, New York, 1958.

The unconditional stability was first discussed by

LLEWELLYN, F. B., Some Fundamental Properties of Transmission Systems. *Proc. IRE* **40**, 271-283 (1952).

The concept of unilateral gain was introduced by

MASON, S. J., Power Gain in Feedback Amplifier. *IRE Trans. CT-1*, 20-25 (1954).

The operating principle of tunnel diodes is based on

ESAKI, L., New Phenomenon in Narrow Germanium $p-n$ Junctions. *Phys. Rev.* **109**, 603-604 (1958).

In order to write the qualitative discussion of semiconductors in Section 7.5, the present author consulted

SHIVE, J. N., "Semiconductor Devices." Van Nostrand, Princeton, New Jersey, 1959.

The Manley-Rowe relation was published by

MANLEY, J. M., and ROWE, H. E., Some General Properties of Nonlinear Elements—I. General Energy Relations. *Proc. IRE* **44**, 904-913 (1956).

The derivation given in Section 7.6 is, however, based on

MANLEY, J. M., and ROWE, H. E., General Energy Relations in Nonlinear Reactances. *Proc. IRE* **47**, 2115-2116 (1959).

The dynamic Q of variable capacitances was first used by

KUROKAWA, K., and UENOHARA, M., Minimum Noise Figure of the Variable Capacitance Amplifier. *Bell System Tech. J.* **40**, 695-722 (1961).

A simple method for measuring the dynamic Q and the adjustment of parametric amplifier circuits was presented by

KUROKAWA, K., Use of Passive Circuit Measurements for Adjustment of Variable Capacitance Amplifiers. *Bell System Tech. J.* **41**, 361–381 (1962).

For the discussion of electron beams, the author consulted Chapter 3 in

SMULLIN, L. D., and HAUS, H. A., "Noise in Electron Devices." Wiley, New York, 1959.

Our n_p and n_s are closely related to Π and S , originally introduced by H. A. Haus:

$$\Pi = n_p/4\pi B, \quad S = n_s/4\pi B$$

The comparison between magnitudes of kinetic and electromagnetic powers is based on

LOUISELL, W. H., and PIERCE, J. R., Power Flow in Electron Beam Devices. *Proc. IRE* **43**, 425–427 (1955).

The distributed couplings of three waves, one circuit, and two space-charge waves was discussed in

PIERCE, J. R., Traveling-Wave Tubes. *Bell System Tech. J.* **29**, 608–669 (1950).

Parameters $4QC$, C , b used in the above paper are given in terms of our parameters by

$$4QC = \frac{\beta_q^2}{(2\beta_q|C_{13}|^2)^{2/3}}, \quad C = \frac{1}{\beta_e} (2\beta_q|C_{13}|^2)^{1/3}, \quad b = \frac{\beta_s - \beta_e}{2(\beta_q|C_{13}|^2)^{1/3}}$$

The derivation of linear differential equations for nonlinear systems, such as oscillators, is known as equivalent linearization. The first comprehensive discussion, which is a translation of Russian monographs published in 1934 and 1937, was given by

KRYLOV, N., and BOGOLIUBOV, N., "Introduction to Nonlinear Mechanics." Princeton Univ. Press, Princeton, New Jersey, 1943.

The discussion of generalized functions presented in Appendix II is based on

LIGHTHILL, M. J., "Introduction to Fourier Analysis and Generalized Functions." Cambridge Univ. Press, London and New York, 1959.

Two good papers on noise in oscillators are

EDSON, W. A., Noise in Oscillators. *Proc. IRE* **48**, 1454–1466 (1960).

MULLEN, J. A., Background Noise in Nonlinear Oscillators. *Proc. IRE* **48**, 1467–1473 (1960).

The discussion in Section 9.3 is based on

KUROKAWA, K., Noise in Synchronized Oscillators. *Trans. IEEE MTT-16*, 234–240 (1968).

The following are general texts covering many branches of the subject:

SOUTHWORTH, G. C., "Principles and Applications of Waveguide Transmission." Van Nostrand, Princeton, New Jersey, 1950.

GINZTON, E. L., "Microwave Measurements." McGraw-Hill, New York, 1957.

COLLIN, R. E., "Field Theory of Guided Waves." McGraw-Hill, New York, 1960.

RAMO, S., and WHINNERY, J. R., "Fields and Waves in Modern Radio." Wiley, New York 1962.

JOHNSON, C. C., "Field and Wave Electromagnetics." McGraw-Hill, New York, 1965.

COLLIN, R. E., "Foundations for Microwave Engineering." McGraw-Hill, New York, 1966.



INDEX

A

Actual power, 36
Adjoint matrix, 204
AM noise, 394
Ampere-turn, 67
Attenuation constant, 23, 162, 198
Autocorrelation function, 423
Available gain, 289
Available power, 34

B

Beam-coupling factor, 361
Bessel function, 141
 hyperbolic, 379
Bilinear transformation, 23
Boltzmann constant, 300
Brillouin zone, 286

C

Canonical form, 296
Canonical transformation, 296
Characteristic admittance, 7
Characteristic impedance, 7, 93
Circulator, 248
Completeness, 85, 107, 110, 120
Conduction band, 322
Conductivity, 69

Conservation

of charge, 67
of energy, 70

Convolution, 421

Coordinate transformation, 205

Coulomb, 67

Coupled mode, 256

Coupling matrix, 258

Cross-correlation function, 424

Cross-power spectral density, 424

Cutoff, 89

Cylindrical coordinate, 61

D

Degeneracy, 140

removal of 148, 160, 167

Delta-function, 415

Depletion layer capacitance, 325

Diagonal matrix, 211

Diffusion capacitance, 326

Directional coupler, 168, 169

Displacement current, 66

Divergence, 48

Dominant mode, 90

Dynamic Q , 340

E

Eigenfunction, 99

Eigenfunction approach, 84
 Eigenvalue, 99, 206
 infinite growth of, 107, 399
 problem, 99, 206
 Eigenvector, 206
 Electron beam, 346
 Ensemble average, 294
 Equivalent circuit
 of electrode gap, 364
 of linear one-port generator, 33
 of lossless reciprocal three-port
 network, 233
 of lossless reciprocal two-port
 network, 33
 of one-port resonant cavity, 187
 of reciprocal two-port network, 31
 of tunnel diode, 330
 of two-port resonant cavity, 191
 of unconditionally stable amplifier, 313
 of waveguide window, 97, 149
 Exchangeable gain, 344
 Exchangeable noise figure, 344
 Exchangeable noise measure, 344
 Exchangeable power, 34
 External Q , 186

F

Faraday rotation, 254
 Fast wave, 353
 Fermi-Dirac distribution law, 321
 Fermi level, 321
 Ferrite, 240
 FM noise, 394
 Foster's reactance theorem, 228
 Fourier expansion, 110
 Fourier transformation, 416
 Functional, 104

G

Gauss's theorem, 49, 51
 Generalized function, 415
 Good function, 415
 Gradient, 54
 Group velocity, 94, 138, 358

H

Helmholtz's theorem, 61
 Hole, 322

I

Idealized model, 1, 4
 Idler, 337
 Incident wave, 12
 Injection locking, 386-389
 Integration by parts, 57
 Inverse matrix, 202
 Isolator, 254

J

Junction capacitance, 326

K

Kinetic energy, 351
 Kinetic power, 350
 Kinetic voltage, 350
 Klystron, 364-367

L

Linearity, 5
 Loaded Q , 190
 Lossless condition, 224, 225, 277, 372
 Low-field loss, 247

M

Magic T, 253
 Magnetic induction, 63
 Magnetization, 242
 Manley-Rowe relations, 334
 Matched load, 98
 Matrix, 200-217
 Maximum-minimum property, 411
 Maxwell's equations, 67, 69
 Movable short, 98

N

n -Dimensional vector, 204
 Natural precession frequency, 246
 Net power, 11

Neumann function, 141
 $1/f$ Noise, 396
 Noise figure, 302
 Noise temperature, 301
 Nonlinear reactance, 333
 Nonlinearity, 333
 Nonsingular matrix, 203
 Nonsingular transformation, 297
 Normalization, 107, 118, 125
 Normalized admittance, 21
 Normalized impedance, 14, 93

O

ω - β Diagram, 286
 Operating noise measure, 301
 Operating noise temperature, 300
 Optimum noise measure, 303
 of electron beam device, 378
 of parametric amplifier, 342
 of tunnel diode amplifier, 331
 Orthogonality, 78, 101, 116, 125
 Orthonormality, 125, 126, 128
 Oscillator, 380
 Overcoupling, 190

P

Parametric amplifier, 337
 Passband, 281
 Perfect conductor, 85
 Periodic structure, 274-282
 Permeability, 69
 Perturbation method, 164, 264
 Phase constant, 6
 Phase velocity, 11, 95, 138
 Photon, 319
 Piecewise-continuous function, 107, 119, 174
 Plane wave, 76
 Plasma frequency, 353
 reduced, 354
 Plasma wavelength, 364
 Poisson's equation, 58
 Polarization, 79-80
 Positive-definiteness, 212
 Power-loss method, 164

Power spectral density, 424
 Power transmission coefficient, 37
 Power wave, 35
 Poynting vector, 70
 complex, 72
 Principle of superposition, 5
 Projection, 42
 Propagation constant, 6
 Pump, 336

Q

Q -Curve, 173

R

Reciprocal condition, 220, 222, 276
 Reference impedance, 222
 Reflected wave, 12
 Reflection coefficient, 12
 power, 37
 power wave, 36
 Rellich's theorem, 401, 407
 Resonant cavity, 170
 Rieke diagram, 385
 Right-handed screw rule, 43
 Rotary vane type attenuator, 147
 Rotation, 51

S

Saturation magnetization, 246
 Scalar products, 42
 Scattering matrix, 222
 Self-adjoint matrix, 204
 Semiconductor, 322-326
 intrinsic, 322
 n -type, 323
 p -type, 323
 Semipositive-definiteness, 215
 Separation of variable, 140
 Similarity transformation, 205
 Simultaneous diagonalization, 215, 216, 253
 Simultaneous matching, 311-312, 345
 Singular matrix, 203
 Shot noise, 330, 397
 Skin depth, 76

Slow wave, 353
Slow-wave structure, 288
Smith chart, 15
Space charge wave, 355
Space harmonics, 287
Spectral representation, 210, 264, 279, 280
Square-integrable function, 107
Standard noise temperature, 302
Standing wave detector, 98
Standing wave method, 21
Standing wave ratio, 20
Stokes's theorem, 52
Stopband, 281
Surface current, 87
Symmetrical Y junction, 234-240

T

Tapered section, 99
Tensor permeability, 241
Total matching, 239
Trace, 207
Transducer gain, 289, 300
Transfer matrix, 276
Transposed matrix, 203
Transverse electric mode, 90
Transverse electromagnetic mode, 123
Transverse magnetic mode, 123
Traveling wave, 11
Traveling wave tube, 367

Tunnel diode, 318-329
Tunnel effect, 328

U

Unconditional stability, 308, 311
Undercoupling, 190
Unilateral gain, 315
Unilateralization, 318
Unitary matrix, 204
Unloaded Q , 186
Up-converter, 336

V

Valence band, 322
Variational expression, 105, 116, 127-129,
145, 152, 167, 169, 173, 175, 196, 197,
208
Vector, 40-63
Vector product, 43

W

Waveguide, 84
impedance, 92
window, 97, 149
Wavelength, 10, 90, 138
Weber, 67
Well-behaved function, 130
White noise, 389

